

John
von
Neumann

COLLECTED
WORKS

Vol. II

Operators
Ergodic Theory
and
Almost
Periodic
Functions



John
von
Neumann

COLLECTED WORKS

Volume II

John von Neumann

COLLECTED WORKS

Volume II

**OPERATORS, ERGODIC THEORY AND
ALMOST PERIODIC FUNCTIONS IN A GROUP**



JOHN VON NEUMANN
1903-1957

John von Neumann

COLLECTED WORKS

GENERAL EDITOR

A. H. TAUB

RESEARCH PROFESSOR OF APPLIED MATHEMATICS
DIGITAL COMPUTER LABORATORY
UNIVERSITY OF ILLINOIS

Volume II

OPERATORS, ERGODIC THEORY AND
ALMOST PERIODIC FUNCTIONS IN A GROUP



PERGAMON PRESS

OXFORD · LONDON · NEW YORK · PARIS

1961

PERGAMON PRESS LTD.,
Headington Hill Hall, Oxford.
4 and 5 Fitzroy Square, London W.1.

PERGAMON PRESS INC.,
122 East 55th Street, New York 22, N.Y.
1404 New York Avenue, N.W., Washington 5, D.C.
Statler Center - 640,
900 Wilshire Blvd., Los Angeles 27, Calif.

PERGAMON PRESS S.A.R.L.
24 Rue des Ecoles, Paris Ve

PERGAMON PRESS G.m.b.H.
Kaiserstrasse 75, Frankfurt-am-Main

This Compilation
Copyright © 1961
Pergamon Press Ltd.

Library of Congress Catalog Card No. 59-14497

Printed in Great Britain by
PERGAMON PRINTING & ART SERVICES LTD.,
LONDON

ACKNOWLEDGEMENT

THE publication of the six volumes of the collected works of John von Neumann represents an imposing burden of work, great and broad knowledge, and time-consuming research. It would hardly have been possible to make this collection available to the scientific community, and surely not in so short a time, had it not been for the welcome fact that Dr. A. H. Taub volunteered to undertake the labor of assembling and editing the manuscript. His dedication and the unselfish and intensive concentration with which he handled the task have made this publication possible. I believe that he was motivated to take on this task by the memory of a close personal friend, a man whom he admired for his scientific work. I would like to be permitted to express my deepest gratitude to A. H. Taub, not only in my own name, but also in the name of the entire scientific community who will benefit by his wonderful work. I know that many colleagues of John von Neumann share my gratitude, and Dr. Oppenheimer of the Institute for Advanced Study, where my late husband spent so many fruitful years, has asked to be associated with this acknowledgement of a great debt to the Editor.

KLARA VON NEUMANN-ECKART

PUBLISHERS' NOTE

The Publishers wish to express their sincere gratitude for the kind cooperation received from the publishers of the various publications in which the articles reproduced in these collected works first appeared, and for permission to reproduce this material. The exact source of each article appears in the Bibliography, listed under the year of publication.

Thanks are due to the following publishers for permission to include material from the journals listed, in the present volume :

To the Springer-Verlag (Berlin, Heidelberg and Göttingen), for the articles which first appeared in the *Mathematische Annalen*.

To Messrs. Walter de Gruyter & Co. (Berlin) for the articles reprinted from the *Journal für die reine und angewandte Mathematik*.

To the Princeton University Press, N.J., for the articles reprinted from *Annals of Mathematics*.

To N.V. Erven P. Noordhoff, Groningen, for the article which appeared in *Compositio Mathematica*.

To the U.S. National Academy of Sciences for the articles reprinted from their *Proceedings*.

To the American Mathematical Society for the articles reprinted from their *Transactions*.

PREFACE

THE following pages contain a reprinting of all the articles published by John von Neumann, some of his reports to government agencies and to other organizations, and reviews of unpublished manuscripts found in his files. The published papers, especially the earlier ones, are given herein in essentially chronological order. Exceptions in this ordering have been made in order that his work in certain fields could be presented as a whole.

A number of reports written by von Neumann could not be included in this collection because they are still classified. Others were not included because they are essentially lecture notes which contain well-known material.

Shortly after von Neumann's death a number of people devoted much time to studying the von Neumann files. One result of this study was the publication of a posthumous paper by von Neumann :

Non-isomorphism of Certain Continuous Rings (with an introduction by I. Kaplansky). *Ann. Math.* 67 (1958), pp. 485-496.

It also became apparent that there were a number of manuscripts which for one reason or another were not suitable for publication but whose existence should be made known to the scientific public, because these manuscripts contained partial results or techniques of wide interest. It was decided to have such manuscripts placed in the library of the Institute for Advanced Study and have reviews of their contents included in this collection of von Neumann's work.*

The Bibliography given at the end of this Volume contains a listing of von Neumann's scientific works, including his books and mimeographed lecture notes. There is noted in this listing the volume and item number where any particular paper or report in the bibliography, included in this collection, may be found. A brief resume of the contents of the various volumes follows.

VOLUME I : Logic, Theory of Sets and Quantum Mechanics—starts with the article entitled, "The Mathematician," first published in 1947, in which von Neumann discusses the nature of the intellectual effort in mathematics. The volume then continues with early papers on the subjects given in the title.

VOLUME II : Operators, Ergodic Theory and Almost Periodic Functions in a Group—includes papers on the subjects listed in the title.

VOLUME III : Rings of Operators—contains his work on, and closely related to, the theory of the rings of operators.

*Additional material, consisting of unfinished manuscripts, notes and some of von Neumann's scientific correspondence, is also on file in this library.

VOLUME IV : Continuous Geometry and other topics—includes the papers on continuous geometry, as well as papers written on a variety of mathematical topics. Also included in this volume are four papers on statistics, and reviews of a number of manuscripts found in his files.

VOLUME V : Design of Computers, Theory of Automata and Numerical Analysis—is devoted to von Neumann's work in these topics.

VOLUME VI : Theory of Games, Astrophysics, Hydrodynamics and Meteorology. Papers on these topics, as well as a number of articles based on speeches delivered by von Neumann, are included in this volume.

The editor acknowledges the debt he owes for the comments made by the people who studied the von Neumann files, and is particularly indebted for the remarks of J. Bigelow, J. G. Charney, C. Eckart, P. Halmos, S. Kakutani, F. J. Murray, E. Teller and E. Wigner.

He gratefully acknowledges the helpful advice given by S. Bochner, K. Godel, Deane Montgomery and S. Ulam. Special thanks are due to the persons who evaluated many manuscripts and reviewed some : J. Bardeen, G. Birkhoff, H. H. Goldstine, I. Halperin, G. A. Hunt, I. Kaplansky, H. W. Kuhn, G. W. Mackey, O. Morgenstern and A. W. Tucker.

The Office of Naval Research gave financial support to the work of organizing the von Neumann files and evaluating manuscripts. Without this support the preparation of this collection would have been seriously hampered.

The editor is grateful to the Institute for Advanced Study for making its facilities available to him and for providing many essential and helpful services. It is his pleasure to acknowledge the debt he owes to Mrs. Elizabeth Gorman, who was formerly secretary to Professor von Neumann, for her extremely valuable and unstintingly given assistance.

A final and most important acknowledgement must be made. This is to the untiring efforts of Klara von Neumann-Eckart, and her help in innumerable ways in making available to the scientific community the various contributions of John von Neumann.

A. H. TAUB

CONTENTS

OF VOLUME II

	Page
ACKNOWLEDGEMENT BY KLARA VON NEUMANN-ECKART ..	v
PREFACE BY PROFESSOR A. H. TAUB	vii
1. ALLGEMEINE EIGENWERTTHEORIE HERMITESCHER FUNKTIONALOPERATOREN .	3
2. ZUR ALGEBRA DER FUNKTIONALOPERATOREN UND THEORIE DER NORMALEN OPERATOREN	86
3. ZUR THEORIE DER UNBESCHRÄNKTEN MATRIZEN	144
4. ÜBER EINEN HILFSSATZ DER VARIATIONSRECHUNG	173
5. ÜBER FUNKTIONEN VON FUNKTIONALOPERATOREN	177
6. ALGEBRAISCHE REPRÄENTANTEN DER FUNKTIONEN " BIS AUF EINE MENGE VOM MASSE NULL "	213
7. DIE EINDEUTIGKEIT DER SCHRÖDINGERSCHEN OPERATOREN	221
8. BEMERKUNGEN ZU DEN AUSFÜHRUNGEN VON HERRN ST. LESNIEWSKI ÜBER MEINE ARBEIT " ZUR HILBERTSCHEN BEWEISTHEORIE "	230
9. DIE FORMALISTISCHE GRUNDLEGUNG DER MATHEMATIK	234
10. ZUM BEWEISE DES MINKOWSKISCHEN SATZES ÜBER LINEARFORMEN	240
11. ÜBER ADJUNGIERTE FUNKTIONALOPERATOREN	242
12. PROOF OF THE QUASI-ERGODIC HYPOTHESIS	261
13. PHYSICAL APPLICATIONS OF THE ERGODIC HYPOTHESIS	274
14. DYNAMICAL SYSTEMS OF CONTINUOUS SPECTRA	278
15. ÜBER EINEN SATZ VON HERRN M. H. STONE	287
16. EINIGE SATZE ÜBER MESSBARE ABBILDUNGEN	294
17. ZUR OPERATORENMETHODE IN DER KLASSISCHEN MECHANIK	307
18. ZUSATZE ZUR ARBEIT " ZUR OPERATORENMETHODE IN DER KLASSISCHEN MECHANIK "	363
19. DIE EINFÜHRUNG ANALYTISCHER PARAMETER IN TOPOLOGISCHEN GRUPPEN	366
20. A KOORDINÁTA-MÉRÉS PONTOSÁGÁNAK HATÁRAI AZ ELEKTRON DIRAC-FÉLE ELMÉLETÉBEN (ÜBER DIE GRENZEN DER KOORDINATENMESSUNGS-GENAUIGKEIT IN DER DIRACSCHEN THEORIE DES ELEKTRONS)	387
21. ON AN ALGEBRAIC GENERALIZATION OF THE QUANTUM MECHANICAL FORMALISM	409
22. ZUM HAARSCHEN MASS IN TOPOLOGISCHEN GRUPPEN	445
23. ALMOST PERIODIC FUNCTIONS IN A GROUP I	454
24. THE DIRAC EQUATION IN PROJECTIVE RELATIVITY	502

CONTENTS OF VOLUME II

25.	ON COMPLETE TOPOLOGICAL SPACES							508
26.	ALMOST PERIODIC FUNCTIONS IN GROUPS II	...		...				528
27.	COMPARISON OF CELLS (REVIEWED BY G. HUNT)							558
	BIBLIOGRAPHY OF JOHN VON NEUMANN	...						559
	COMPLETE CONTENTS OF VOLUMES I-VI							567

ERRATA for

Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren

p. 50, line -2 should read: ist Teil von $\mathfrak{M}_q$

p. 62, line 5 should read: . . . so dass stets $(f, Sf) \geq C \cdot |f|^2$. . .

p. 68, line 12 should read: Beweis. . . da $f_a(x)g_a(x) = 0$, $f_a(x)$

line 13 should read: $-g_a(x) = x - a$. . .

line -11 should read: Beweis. . . $f_b(x) \geq 0$, $f_a(x) - f_b(x) \geq 0$

line -10 should read: ist, also $f_b(R)$, $f_a(R) - f_b(R)$ definit. . .

p. 69, line 10 should read: . . ., und $\mathfrak{N}(R, a) = \xi - \mathfrak{M}(R, a)$. . .

line 18 should read: $E(A, a, b) = P_{\mathfrak{M}(R, a)} \cdot P_{\mathfrak{M}(S, b)}$

line 2 of footnote should read:
$$= \sum_{v=1}^p (RP_{\mathfrak{M}(R, x_v)} \cdot \mathfrak{N}(R, x_{v-1})f, f) =$$
$$\sum_{v=1}^p (RP_{\mathfrak{M}(R, x_i)} \cdot \mathfrak{N}(R, x_{i-1})f, P_{\mathfrak{M}(R, x_i)} \cdot \mathfrak{N}(R, x_{i-1})f).$$

line 3 of footnote should read:

. . . mit $P_{\mathfrak{M}(R, x_v)} - P_{\mathfrak{M}(R, x_{v-1})} = P_{\mathfrak{M}(R, x_i)} \cdot \mathfrak{N}(R, x_{i-1})$. . .

line 4 of footnote should read:

also $P_{\mathfrak{M}(R, x_v)} \cdot \mathfrak{N}(R, x_{v-1}) RP_{\mathfrak{M}(R, x_i)} \cdot \mathfrak{N}(R, x_{i-1}) =$

$RP_{\mathfrak{M}(R, x_v)} \cdot \mathfrak{N}(R, x_{i-1})$ gilt.) Nach Satz 7*

p. 71, lines 2 and 3 should read:

$$= \mathfrak{B}_{\mathfrak{M}(R, a'')} \mathfrak{B}_{\mathfrak{M}(S, b'')} - \mathfrak{B}_{\mathfrak{M}(R, a'')} \mathfrak{B}_{\mathfrak{M}(S, b')} - \mathfrak{B}_{\mathfrak{M}(Ra')} \mathfrak{B}_{\mathfrak{M}(S, b'')} \\ + \mathfrak{B}_{\mathfrak{M}(Ra')} \mathfrak{B}_{\mathfrak{M}(S, b')}$$

$$= \mathfrak{B}_{\mathfrak{M}(Ra'')} \mathfrak{B}_{\mathfrak{M}(Sb'')} - \mathfrak{M}(S, b') - \mathfrak{B}_{\mathfrak{M}(Ra')} \mathfrak{B}_{\mathfrak{M}(S, b'')} - \mathfrak{M}(S, b')$$

$$= \mathfrak{B}_{\mathfrak{M}(R, a'')} [\mathfrak{M}(S, b'') - \mathfrak{M}(S, b')] - \mathfrak{B}_{\mathfrak{M}(R, a')} [\mathfrak{M}(S, b'') - \mathfrak{M}(S, b')]$$

$$= \mathfrak{B}_{\mathfrak{M}(R, a'')} \mathfrak{N}(S, b') \mathfrak{M}(Sb'') - \mathfrak{B}_{\mathfrak{M}(Ra')} \mathfrak{N}(Sb') \mathfrak{M}(Sb'')$$

$$= \mathfrak{B}_{\mathfrak{M}(R, a'')} \mathfrak{N}(Ra') \mathfrak{M}(Sb'') \mathfrak{N}(Sb')$$

line 1 of footnote 69 should read:

Da f in $\mathfrak{M}(R, a'')$ und $\mathfrak{N}(R, a')$ liegt, . . .

line 2 of footnote 70 should read:

. . . Für $a'' \leq c'$ ist $\mathfrak{M}(R, c')$

line 3 of footnote 70 should read:

Teil von $\mathfrak{M}(R, a'')$, also $\mathfrak{M}(R, a'') \perp \mathfrak{N}(R, c')$

p. 73, line 10 should read:

. . . so dass stets $F(\lambda_n)f \rightarrow Ef$ ist;

Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren

Einleitung.

I. Wenn wir alle Funktionen eines gewissen metrischen Raumes Ω (z. B. der reellen Zahlengeraden, der komplexen Zahlenebene, der Oberfläche der Einheitskugel, der Strecke 0, 1 usw.) betrachten, die gewissen Regularitätsbedingungen genügen (z. B. stetig und bis auf endlich viele Knicke stetig differenzierbar sind, zweimal stetig differenzierbar sind, ein endliches Absolutwertquadratintegral über Ω haben¹⁾ usw.) und eventuell auch noch gewissen Randbedingungen unterworfen sind, so wird unter einem linearen Operator bekanntlich das Folgende verstanden: Eine Zuordnung R , die jeder Funktion f unserer Klasse eine Funktion Rf (die nicht mehr zu dieser Klasse gehören muß) zuordnet, aber so, daß stets

$$R(a_1 f_1 + \dots + a_k f_k) = a_1 Rf_1 + \dots + a_k Rf_k \quad (a_1, \dots, a_k \text{ komplexe Konstante})$$

ist.

Wenn in Ω ein allgemeiner Maßbegriff (etwa im Sinne des Lebesgueschen) existiert und dv das Volumelement in Ω ist (auf der Geraden: dx , in der Ebene: $dx dy$, auf der Oberfläche der Einheitskugel: $\sin \vartheta \cdot d\vartheta d\varphi$ usw.), und das Integral über diesen ganzen Raum mit $\int_{\Omega}$ bezeichnet wird, so heißt ein linearer Operator R selbstadjungiert oder Hermitesch, wenn für alle zugelassenen Funktionen f, g

$$\int_{\Omega} f \cdot \overline{Rg} \cdot dv = \int_{\Omega} Rf \cdot \bar{g} \cdot dv$$

gilt²⁾).

¹⁾ Die Funktionen dürfen komplexe Werte haben.

²⁾ $\bar{f}$ ist diejenige Funktion, die überall den zum Werte von f konjugiert-komplexen Wert annimmt. — Unsere Terminologie weicht von der üblichen etwas ab:

(Fortsetzung der Fußnote ²⁾ auf nächster Seite.)

Hermitesche lineare Operatoren (wir wollen sie kurz H. O. nennen) sind z. B. die folgenden (wir geben nun auch die Argumente der Funktionen an, und zwar sei P der allgemeine Punkt von Ω):

$$Rf(P) = f(P)$$

$$Rf(P) = x \cdot f(P) \quad (x \text{ irgendeine Koordinate von } P \text{ in } \Omega),$$

$$Rf(P) = \int_{\Omega} \varphi(P', P) \cdot f(P') \cdot dv' \quad (\varphi(P', P) = \overline{\varphi(P, P')}, \text{ d. h. der Kern Hermitesch}),$$

$$Rf(P) = Lf(P) \quad (L \text{ irgendein selbstadjungierter Differential- oder partieller Differentialausdruck, vgl. Anm. 2)).$$

Bekanntlich bestimmen die beiden letzten Typen von H. O. ein Eigenwertproblem, das bei hinreichender Regularität des Kernes $\varphi(P, P')$ bzw. der Koeffizienten von L (nämlich wenn nur ein „Punktspektrum“ existiert) so formuliert werden kann³⁾:

Ein Eigenwert ist eine Zahl λ , zu der es eine Funktion $f \neq 0$ mit

$$Rf = \lambda f$$

gibt; f ist dann Eigenfunktion. Es gibt im allgemeinen unendlich viele Eigenwerte $\lambda_1, \lambda_2, \dots$, denen Eigenfunktionen $\varphi_1, \varphi_2, \dots$ so zugeordnet werden können, daß sie ein vollständiges System von Orthogonalfunktionen bilden. D. h. es gilt:

$$\int_{\Omega} \varphi_m \cdot \overline{\varphi_n} \cdot dv = \begin{cases} 0, & \text{für } m \neq n \\ 1, & \text{für } m = n \end{cases} \quad (\text{Orthogonalität}),$$

ein Differentialoperator R z. B. heißt ja gewöhnlich dann selbstadjungiert, wenn $f \cdot \overline{Rg} - Rf \cdot \overline{g}$ vollständiges Differential eines Ausdruckes der f, g und ihrer Ableitungen ist, also

$$\int_{\Omega} f \cdot \overline{Rg} \cdot dv - \int_{\Omega} Rf \cdot \overline{g} \cdot dv$$

ein nur von den Randwerten von f, g und ihrer Ableitungen abhängiger Ausdruck. Unsere Selbstadjungiertheit (Hermitescher Charakter) folgt hieraus nur, wenn dieser Ausdruck, infolge geeigneter Randbedingungen, verschwindet. D. h. ein im gewöhnlichen Sinne selbstadjungierter Differentialoperator ist es für uns eventuell erst in einem hinreichend eingeeengten Definitionsbereiche (vgl. auch Einleitung VII).

³⁾ Das Eigenwertproblem der regulären Integraloperatoren wurde von Hilbert gelöst (Göttinger Nachr. 1904, S. 49—91), seine Methode wurde später von Courant, Hellinger, F. Riesz, E. Schmidt, Toeplitz u. a. bedeutend vereinfacht. Die Differentialoperatoren (zumindest die selbstadjungierten von zweitem Grade) können bei der Lösung des Eigenwertproblems durch Integraloperatoren ersetzt werden (Hilbert, Göttinger Nachr. 1904, S. 213—259), die dieselben Eigenfunktionen, aber reziproke Eigenwerte haben.

und für jedes Paar zugelassener Funktionen f, g

$$\int_{\Omega} f \cdot \bar{g} \cdot dv = \sum_{n=1}^{\infty} \left(\int_{\Omega} f \cdot \overline{\varphi_n} \cdot dv \right) \cdot \overline{\left(\int_{\Omega} g \cdot \overline{\varphi_n} \cdot dv \right)} \quad (\text{Vollständigkeit}).$$

Wie für alle vollständigen Orthogonalsysteme, gilt der Entwicklungssatz mit Mittelkonvergenz:

$$\int_{\Omega} \left| f - \sum_{n=1}^N a_n \cdot \varphi_n \right|^2 \cdot dv \rightarrow 0 \quad (a_n = \int_{\Omega} f \cdot \overline{\varphi_n} \cdot dv)$$

für $N \rightarrow \infty$.

Wenn sich $\varphi(P, P')$ bzw. die Koeffizienten von L nicht mehr ganz regulär verhalten, so treten Verwicklungen auf, die man durch das Auftreten eines Streckenspektrums kennzeichnet. Immerhin können die Begriffe Eigenwert und Eigenfunktion durch geeignete Verallgemeinerungen (insbesondere durch Herabsetzung der Regularitätsanforderungen an die Eigenfunktion oder gar durch Betrachtung sogenannter differentieller Eigenfunktionen⁴⁾) aufrechterhalten werden. Auch die Vollständigkeitssätze bleiben bei geeigneter Verallgemeinerung (Ersetzen der Summe durch ein Integral usw.) gewahrt; aber der ganze Aufbau wird wesentlich komplizierter und verliert viel von seiner ursprünglichen Anschaulichkeit⁵⁾.

II. Wenn wir die in I. betrachteten kontinuierlichen Räume Ω durch den „diskreten Raum“ $1, 2, 3, \dots$ der positiven ganzen Zahlen ersetzen, wobei die Funktionen $f(P)$ (P durchläuft Ω) durch die Folgen a_n ($n = 1, 2, \dots$) und die Integrale $\int_{\Omega} f \cdot dv$ durch die Summen $\sum_{n=1}^{\infty} a_n$ sinngemäß zu ersetzen sind, so gelangen wir zu ganz analogen Begriffsbildungen wie in I.

Die einfachsten Beispiele von H. O. (ja, wie man nach präziserer Fassung des Begriffes zeigen kann, die einzigen) sind die linearen Transformationen (unendlich vieler Variablen) mit Hermitescher Matrix:

$$R(x_1, x_2, \dots) = (y_1, y_2, \dots),$$

$$y_n = \sum_{m=1}^{\infty} \alpha_{nm} x_m \quad (n = 1, 2, \dots), \quad \alpha_{nm} = \overline{\alpha_{mn}}.$$

(Die Rolle der „Regularitätsbedingungen“ bei Funktionen in Ω spielen da

⁴⁾ Hellinger, Göttinger Inaugural-Dissertation 1907, S. 7; Carleman, Sur les équations intégrales à noyau réel et symétrique, Upsala 1923, S. 82.

⁵⁾ Das Eigenwertproblem der Integralgleichungen mit „beschränktem“ Kerne wurde von Weyl gelöst (Göttinger Inaugural-Dissertation 1908) mit Hilfe der Hilbertschen Theorie der beschränkten Bilinearformen (Göttinger Nachr. 1906, S. 157—209). Auch für eine ausgedehnte Klasse selbstadjungierter Differentialoperatoren zweiten Grades, deren Koeffizienten singularär sein dürfen, erledigte Weyl das Eigenwertproblem (Math. Annalen 68 (1910), S. 220—269). Eine noch allgemeinere Klasse von Integraloperatoren untersuchte Carleman (vgl. Anm. 4)).

Bedingungen von der Art, daß alle $\sum_{m=1}^{\infty} \alpha_{mn} x_m$ konvergieren sollen, u. ä. Wir werden diese Matrizen noch im Anhang III eingehender betrachten.) Diese linearen Transformationen sind offenbar Analoga der Integralkern-Transformationen in kontinuierlichen Räumen Ω :

$$Rf(P) = \int_{\Omega} \varphi(P', P) \cdot f(P') \cdot dv'.$$

Auch für diese Transformationen bzw. für die zu ihnen gehörigen Hermiteschen Formen

$$\sum_{m, n=1}^{\infty} \alpha_{mn} x_m \overline{y_n}$$

ist das Eigenwertproblem von Hilbert gestellt und gelöst worden, falls die Matrix $\{\alpha_{mn}\}$ (bzw. die Hermitesche Form $\sum_{m, n=1}^{\infty} \alpha_{mn} x_m \overline{y_n}$) gewissen Bedingungen genügt, die der „Regularität“ des Kernes $\varphi(P, P')$ bei Integralkern-Transformationen entsprechen.

Es treten dieselben Erscheinungen auf: Bei hoher Regularität (Vollstetigkeit) bloß Punktspektrum, mit denselben Eigenschaften wie in I., nur daß $\int_{\Omega}$ durch $\sum_{n=1}^{\infty}$ zu ersetzen ist. Also: λ ist Eigenwert, wenn es eine Folge $(x_1, x_2, \dots)$ mit

$$R(x_1, x_2, \dots) = \lambda(x_1, x_2, \dots), \text{ d. h. } \sum_{m=1}^{\infty} \alpha_{mn} x_m = \lambda x_n \quad (n = 1, 2, \dots)$$

gibt [ohne daß $(x_1, x_2, \dots) = (0, 0, \dots)$ ist, außerdem ist als „Regularitätsbedingung“ die Endlichkeit von $\sum_{n=1}^{\infty} |x_n|^2$ zu fordern], und $(x_1, x_2, \dots)$ ist dann Eigenfolge. Sind $\lambda_1, \lambda_2, \dots$ alle Eigenwerte, so können ihnen Eigenfolgen $(x_{11}, x_{12}, \dots)$, $(x_{21}, x_{22}, \dots)$, ... so zugeordnet werden, daß sie ein vollständiges Orthogonalsystem bilden:

$$\sum_{m=1}^{\infty} x_{pm} \overline{x_{qm}} = \begin{cases} 1, & \text{für } p = q \\ 0, & \text{für } p \neq q \end{cases} \quad (\text{Orthogonalität}),$$

und⁶⁾

$$\sum_{m=1}^{\infty} x_{mp} \overline{x_{mq}} = \begin{cases} 1, & \text{für } p = q \\ 0, & \text{für } p \neq q \end{cases} \quad (\text{Vollständigkeit}).$$

Bei geringerer Regularität (Beschränktheit) kann wieder ein Strecken-

⁶⁾ Dies ist bekanntlich dem eigentlichen Analogon der Vollständigkeitsrelation (vgl. I.)

$$\sum_{n=1}^{\infty} a_n \overline{b_n} = \sum_{p=1}^{\infty} \left(\sum_{n=1}^{\infty} a_n \overline{x_{pn}} \right) \overline{\left(\sum_{n=1}^{\infty} b_n \overline{x_{pn}} \right)}$$

gleichbedeutend.

spektrum auftreten, was dieselben Komplikationen bedingt, wie die unter I. für kontinuierliche Räume skizzierten; da wir uns später noch recht eingehend damit zu beschäftigen haben werden, wollen wir es hier nicht weiter erörtern⁷⁾. Irgendwelche „Regularitäts“-Annahmen wurden übrigens in der bisherigen Literatur (sowohl bei Matrizen als auch bei Integraloperatoren) stets gemacht: eine Diskussion des Eigenwertproblems auf Grund des Hermiteschen Charakters allein erfolgte nie.

III. Das analoge Verhalten der kontinuierlichen Räume Ω untereinander, sowie mit dem hier betrachteten diskreten Raume 1, 2, ..., liegt bekanntlich an folgendem:

Wenn $\varphi_1, \varphi_2, \dots$ ein vollständiges System von Orthogonalfunktionen im betrachteten Raume Ω ist, so wird durch die Zuordnung

$$f \leftrightarrow (x_1, x_2, \dots), \quad x_n = \int_{\Omega} f \cdot \overline{\varphi_n} \cdot dv \quad (n = 1, 2, \dots)$$

jedem f mit endlichem Absolutwertquadratintegral $\int_{\Omega} |f|^2 \cdot dv$ eine Folge $(x_1, x_2, \dots)$ mit endlicher Absolutwertquadratsumme $\sum_{n=1}^{\infty} |x_n|^2$ zugeordnet. Dabei entsprechen verschiedenen f verschiedene $(x_1, x_2, \dots)$, und nach dem Satz von Fischer und F. Riesz⁸⁾ entspricht wirklich jedes $(x_1, x_2, \dots)$ mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$ einem f mit endlichem $\int_{\Omega} |f|^2 \cdot dv$.

Die Zuordnung ist linear, d. h.

$$\text{aus } f \leftrightarrow (x_1, x_2, \dots) \quad \text{folgt} \quad af \leftrightarrow (ax_1, ax_2, \dots),$$

$$\text{aus } \left\{ \begin{array}{l} f \leftrightarrow (x_1, x_2, \dots) \\ g \leftrightarrow (y_1, y_2, \dots) \end{array} \right\} \quad \text{folgt} \quad f + g \leftrightarrow (x_1 + y_1, x_2 + y_2, \dots),$$

und es ist

$$\int_{\Omega} f \cdot \bar{g} \cdot dv = \sum_{n=1}^{\infty} x_n \overline{y_n} \quad ^9).$$

D. h. alle bei der Beschreibung der Eigenwerte und Eigenfunktionen verwendeten Begriffsbildungen sind bei dieser Zuordnung invariant; also müssen dieselben in den in I. betrachteten kontinuierlichen Räumen und im diskreten Raume aus II. dasselbe Verhalten zeigen.

⁷⁾ Die vollstetigen und die beschränkten Formen (bzw. Operatoren) wurden von Hilbert entdeckt, und er löste ihr Eigenwertproblem (Göttinger Nachr. 1906, S. 157—209), weiter untersucht wurden die beschränkten Bilinearformen von Hellinger, vgl. Anm. ⁴⁾, sowie Journal f. Math. **136** (1909), S. 210—273.

⁸⁾ Z. B. Göttinger Nachr. 1907, S. 116—122.

⁹⁾ Den Beweis dieser, allgemein bekannten, Tatsachen (zusammen mit dem Fischer-F. Riesz'schen Satze) werden wir im Kap. I und im Anhang I erbringen.

Diese Überlegungen sind es bekanntlich im wesentlichen, durch welche Hilbert die Eigenwerttheorie der Integraloperatoren in kontinuierlichen Räumen auf die der H. O. für Folgen $(x_1, x_2, \dots)$ (d. h. Funktionen des diskreten Raumes $1, 2, \dots$) zurückgeführt hat¹⁰⁾.

IV. Der Gegenstand dieser Arbeit ist die Aufstellung einer allgemeinen Eigenwerttheorie *aller* H. O. Dabei werden wir abstrakt vorgehen, d. h. es von vornherein so einrichten, daß sich unsere Resultate gleichmäßig auf alle in I. genannten kontinuierlichen Räume anwenden lassen, sowie auf den diskreten Raum in II. (d. h. auf die Operatoren der Funktionen in denselben).

Zu diesem Zwecke werden wir im Kap. I eine allgemeine Charakterisierung (durch fünf Eigenschaften A bis E, vgl. dort) angeben, die auf einen jeden der folgenden Funktionenräume paßt:

Alle (komplexen, im Lebesgueschen Sinne meßbaren) Funktionen f in einem Raume Ω (Zahlengerade, Zahlenebene, Oberfläche der Einheitskugel, Intervall $0, 1$ usw., vgl. Anhang I) mit dem Volumelement dv , für die $\int_{\Omega} |f|^2 \cdot dv$ endlich ist. Alle Folgen $(x_1, x_2, \dots)$ von (komplexen) Zahlen mit endlichem

$$\sum_{n=1}^{\infty} |x_n|^2.$$

(Zum ersten Falle ist zu bemerken, daß solche Funktionen f, g , für die $f(P) = g(P)$ überall — mit Ausnahme einer Lebesgueschen 0-Menge — gilt, als nicht verschieden anzusehen sind.)

Die Elemente dieser Funktionenräume lassen alle Vektor-Operatoren zu: man kann sie addieren und mit (komplexen) Zahlen multiplizieren, ferner ist für zwei von ihnen das „innere Produkt“ $\int_{\Omega} f \cdot \bar{g} \cdot dv$ bzw. $\sum_{n=1}^{\infty} x_n \bar{y}_n$ definiert. Das innere Produkt bezeichnen wir mit (f, g) (wenn nichts Besonderes gesagt wird, wollen wir auch die Folgen $(x_1, x_2, \dots), (y_1, y_2, \dots), \dots$ als Funktionen ansehen — vgl. II. — und mit $f, g, \dots$ bezeichnen), mit seiner Hilfe ist in allen Fällen der Hermitesche Charakter der Operatoren R zu definieren:

$$(f, Rg) = (Rf, g).$$

(In I. war das die Definition, und der in II. ist es gleichwertig.) Wie bei Vektoren ermöglicht das innere Produkt die Definition eines Absolutwertes: $|f| = \sqrt{(f, f)} \left(\sqrt{\int_{\Omega} |f|^2 \cdot dv} \text{ bzw. } \sqrt{\sum_{n=1}^{\infty} |a_n|^2} \right)$, und einer Entfernung:

¹⁰⁾ Dies ist insofern historisch ungenau, als der Fischer-F. Rieszsche Satz erst nachher gefunden wurde. Dieselben Überlegungen haben in der neuen Quantentheorie eine wichtige Rolle gespielt, vgl. Schrödinger, Annalen d. Phys. 79, 8 (1926), S. 734—756.

$|f - g|$ (diese ist also das Analogon der gewöhnlichen Euklidischen Entfernung); unser Raum wird demnach durch dasselbe „metrisiert“ und „topologisiert“. D. h. topologische Ausdrucksweisen, wie Abgeschlossenheit, Stetigkeit usw., werden in ihm sinnvoll.

Es ist vielleicht nicht überflüssig, noch ein Wort der konsequenten Beschränkung auf Funktionen f mit endlichem $\int_{\Omega} |f|^2 \cdot dv$ zu widmen. Für Eigenwertprobleme bei Folgen ist es seit den grundlegenden Untersuchungen von Hilbert (vgl. Anm. 7)) und E. Schmidt¹¹⁾ üblich, die Endlichkeit von $\sum_{n=1}^{\infty} |x_n|^2$ zu verlangen; bei den f in kontinuierlichen Räumen ist jedenfalls für die Eigenfunktionen des Punktspektrums $\int_{\Omega} |f|^2 \cdot dv$ endlich. Im Streckenspektrum ist freilich weder $\sum_{n=1}^{\infty} |x_n|^2$ noch $\int_{\Omega} |f|^2 \cdot dv$ endlich (wenn überhaupt noch Eigenfolgen bzw. -funktionen existieren und nicht zu „differentiellen Lösungen“ Zuflucht genommen werden muß), indessen werden wir (unter Verallgemeinerung der Hilbertschen Methoden) zeigen, daß es gerade günstig ist, die Eigenfunktionen im Streckenspektrum als uneigentliche Gebilde zu behandeln, d. h. das Eigenwertproblem ohne ihre explizite Nennung zu formulieren.

Unser Funktionenraum ist, wenn wir den diskreten Raum $1, 2, \dots$ zugrunde legen, der sogenannte (komplexe) Hilbertsche (d. h. unendlichvieldeimensionale Euklidische) Raum: die Menge aller Folgen $(x_1, x_2, \dots)$ mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$. Da, wie in III. bemerkt wurde, alle unsere Funktionenmannigfaltigkeiten isomorph sind (d. h. ein-eindeutig aufeinander abgebildet werden können, wobei die Operationen $f + g, a \cdot f, (f, g)$ in ihre Analoga übergehen), wollen wir sie auch als Hilbertsche Räume bezeichnen. Und sie alle sollen als spezielle Verwirklichungen des allgemeinen abstrakten Hilbertschen Raumes angesehen werden, der durch die „inneren“ Eigenschaften A bis E allein (und bis auf Isomorphismen vollständig, vgl. Kap. I) charakterisiert ist.

Den abstrakten Hilbertschen Raum nennen wir $\S$.

V. Im abstrakten Hilbertschen Raume verstehen wir unter einem H. O., wie in I. gesagt wurde, eine Funktion R (für gewisse f aus $\S$ definiert und Werte aus $\S$ annehmend), die die folgenden Eigenschaften besitzt:

- (α)
$$R(a_1 f_1 + \dots + a_k f_k) = a_1 R f_1 + \dots + a_k R f_k$$

$$(a_1, \dots, a_k \text{ komplexe Konstanten}),$$
- (β)
$$(f, R g) = (R f, g)$$

¹¹⁾ Z. B. Rend. d. Circ. Mat. d. Palermo 25 (1908), S. 57—73.

(soweit beide Seiten Sinn haben). Wir nehmen übrigens einstweilen an, daß der Definitionsbereich von R eine lineare Mannigfaltigkeit ist, d. h. mit $f_1, \dots, f_k$ auch $a_1 f_1 + \dots + a_k f_k$ enthält. (Die endgültige Definition der H. O. wird zwar etwas anders lauten, es kommt uns aber zunächst darauf nicht an.)

Wir nennen einen H. O. R beschränkt, wenn $|f| \leq 1 \quad |Rf| \leq C$ zur Folge hat, wobei C beliebig, aber von f unabhängig ist. Da allgemein $|af| = |a| |f|$ gilt, ist dies mit $|Rf| \leq C \cdot |f|$, oder auch $|Rf - Rg| \leq C \cdot |f - g|$ gleichbedeutend: d. h. es drückt eine Art Stetigkeit von R aus¹²⁾.

In der von Hilbert aufgebauten Theorie der beschränkten Hermiteschen Formen (d. h. H. O.) kann man stets annehmen, daß R für alle f von $\mathfrak{H}$ definiert ist: wenn das nicht der Fall ist, kann es (wegen seiner Stetigkeit) auf ganz $\mathfrak{H}$ fortgesetzt werden. Wenn wir aber eine allgemeine Theorie der H. O. aufstellen wollen, die auch über die beschränkten hinausgeht, so dürfen wir nicht verlangen, daß R überall in $\mathfrak{H}$ definiert sei. (Nach einem Satze von Toeplitz muß sogar jeder überall sinnvolle H. O. beschränkt sein: vgl. Satz 48, α), γ) und Anm. ⁶¹⁾.) So sei z. B. Ω das Intervall $0, 1$, und R der Operator $i \frac{d}{dx} \dots$ für alle differentiierbaren, in 0 und 1 verschwindenden Funktionen. R ist ein H. O., aber es ist offenbar nicht für undifferentiierbare Funktionen zu definieren! Ebenso wenig wird für einen hochsingulären Integralkern $\varphi(P, P') \int_{\Omega} \varphi(P', P) \cdot f(P') \cdot dv'$ für alle f (mit endlichem $\int_{\Omega} |f|^2 \cdot dv$) konvergieren. Ein anders geartetes, aber auch recht charakteristisches Beispiel ist dieses: Ω sei die Zahlengerade, $Rf(x) = x \cdot f(x)$. Auch R ist ein H. O., da aber $\int_{-\infty}^{\infty} |f(x)|^2 dx$ endlich sein kann bei unendlichem $\int_{-\infty}^{\infty} x^2 |f(x)|^2 dx$, ist auch dieses R nicht stets sinnvoll.

Die Einschränkung also, daß (α) , (β) nur soweit zu gelten haben, als beide Seiten sinnvoll sind, ist eine recht wesentliche. Immerhin werden wir aber (als Minimum) verlangen, daß

(γ) der Definitionsbereich von R überall dicht (in $\mathfrak{H}$) sei.

Was dies bedeutet, kann man sich z. B. an $i \frac{d}{dx} \dots$ bzw. $x \dots$ mühelos

¹²⁾ Wenn wir $\mathfrak{H}$ als Folgenraum (eigentlichen Hilbertschen Raum) auffassen, so ist unsere Definition der Beschränktheit dem Wortlaute nach von der Hilbertschen verschieden: wir verlangen (in Hilberts Terminologie) die Beschränktheit des mit sich selbst gefalteten R . Immerhin sind diese Formulierungen gleichwertig, näher erörtern wir dies in Kap. II, Satz 12. — Man beachte übrigens, daß unsere ganze Topologie von $\mathfrak{H}$ stets im Sinne der sogenannten „starken Konvergenz“ zu verstehen ist, vgl. Weyl, Dissertation, S. 9.

klarmachen: diese sind für alle (in 0 und 1 verschwindenden) Polynome bzw. für alle außerhalb eines willkürlichen (aber endlichen) Intervalles stets verschwindenden Funktionen immer sinnvoll, und beide Funktionen-Mengen sind überall dicht.

Bei dieser Anordnung laufen wir aber Gefahr, zu wenig zu verlangen: unsere H. O. könnten bereits an solchen Stellen (infolge einer willkürlichen Einengung des Definitionsbereiches) sinnlos sein, wo sie noch auf natürliche Weise definiert werden könnten. (So hätte man z. B. bei $i \frac{d}{dx} \dots$, unbeschadet der überall-Dichtigkeit des Definitionsbereiches, statt einmaliger Differenzierbarkeit auch Analytizität verlangen können.) Um dieser Schwierigkeit zu entgehen, stellen wir den Begriff des maximalen H. O. auf:

Der H. O. S heiße Fortsetzung des H. O. R , wenn überall wo Rf Sinn hat, auch Sf Sinn hat, und mit Rf übereinstimmt. Wenn $S \neq R$ ist (d. h. S an mehr Stellen sinnvoll ist, wie R), so ist S eine eigentliche Fortsetzung. Ein H. O. R , zu dem es keine eigentlichen Fortsetzungen mehr gibt, heiße maximal. (Abkürzungen: Forts., max.)

Da, wie wir sehen werden (Satz 35), jeder H. O. zu einem max. H. O. fortsetzbar ist, müssen wir in erster Linie die max. H. O. untersuchen. Dazu ist es aber notwendig, eine allgemeine Formulierung des Eigenwertproblems anzugeben (und zwar im abstrakten Hilbertschen Raume) bei der Punkt- und Streckenspektrum gleichmäßig Berücksichtigung finden. Dies soll im nächsten Paragraphen, in Anlehnung an Hilberts Formulierung des Problems, erfolgen.

VI. Betrachten wir zunächst den k -dimensionalen Euklidischen Raum. Wenn dort ein H. O. gegeben ist, d. h. eine Zuordnung

$$R(x_1, \dots, x_k) = (y_1, \dots, y_k),$$

$$y_n = \sum_{m=1}^k \alpha_{mn} x_m \quad (n = 1, \dots, k), \quad \alpha_{mn} = \overline{\alpha_{nm}}$$

(Konvergenzfragen gibt es hier ja nicht!), so ist es bekanntlich stets möglich R durch Einführung neuer kartesischer Koordinaten (d. h. durch eine unitäre Transformation) auf die Diagonalform zu bringen. D. h. es gibt eine unitäre Transformationsmatrix $\{u_{mn}\}$

$$\sum_{p=1}^k u_{mp} \overline{u_{np}} = \begin{cases} 1, & \text{für } m = n \\ 0, & \text{für } m \neq n \end{cases}, \quad \sum_{p=1}^k u_{pm} \overline{u_{pn}} = \begin{cases} 1, & \text{für } m = n \\ 0, & \text{für } m \neq n \end{cases},$$

so daß nach der Transformation

$$\xi_n = \sum_{m=1}^k u_{mn} x_m, \quad \eta_n = \sum_{m=1}^k u_{mn} y_m \quad (n = 1, \dots, k)$$

R einfach durch

$$\eta_n = \lambda_n \xi_n \quad (n = 1, \dots, k)$$

($\lambda_1, \dots, \lambda_k$ reelle Konstante, die Eigenwerte von R !) beschrieben wird¹³). D. h. es ist

$$\alpha_{mn} = \sum_{p=1}^k \lambda_p u_{mp} \overline{u_{np}} \quad (m, n = 1, \dots, k).$$

Dabei sind die $\lambda_1, \dots, \lambda_k$, bis auf die Reihenfolge, durch R eindeutig bestimmt, nicht aber die u_{mn} ! Wie man sofort sieht, kann man sie durch $\vartheta_n u_{mn}$ ersetzen ($\vartheta_1, \dots, \vartheta_k$ beliebige Zahlen vom Absolutwerte 1), und wenn mehrere der λ_n einander gleich sind, sind sogar die dazugehörigen Zeilen ($u_{1n}, \dots, u_{kn}$) beliebig untereinander unitär transformierbar.

Eine uneindeutige Formulierung des Eigenwertproblems, wie diese (Transformierbarkeit auf die Diagonalforn), ist wenig geeignet, um von den k -dimensionalen Euklidischen Räumen auf den unendlichvioldimensionalen (d. h. Hilbertschen) übertragen zu werden. Die richtige, leicht zu verallgemeinernde, Fassung des endlichvioldimensionalen Problems findet man, indem man es nach Hilbert eindeutig formuliert. Dies geschieht so:

Wir setzen

$$e_{mn}(\lambda) = \sum_{\lambda_p \leq \lambda} u_{mp} \overline{u_{np}} \quad (\lambda \text{ eine beliebige reelle Zahl}),$$

dann sind die $\{e_{mn}(\lambda)\}$ Matrizen von H. O. $E(\lambda)$, mit der leicht zu verifizierenden Eigenschaft

$$(\alpha) \quad E(\lambda) E(\mu) = E(\mu) E(\lambda) = E(\lambda) \quad (\lambda \leq \mu)$$

(also insbesondere $E(\lambda)^2 = E(\lambda)$)¹⁴). Seien $l_1, \dots, l_h$ ($h \leq k$) die voneinander verschiedenen Werte der $\lambda_1, \dots, \lambda_k$, und zwar wachsend angeordnet. Dann ist $E(\lambda)$ (als Funktion von λ) in den Intervallen

$$\lambda < l_1, \quad l_1 \leq \lambda < l_2, \quad \dots, \quad l_{h-1} \leq \lambda < l_h, \quad l_h \leq \lambda$$

konstant, und bei $l_1, \dots, l_h$ liegen (links) Sprünge. Im ersten der genannten Intervalle ist es offenbar $= 0$, im letzten (wegen der Unitarität von $\{u_{mn}\}$) hat es die Einheitsmatrix, es ist also der identische Operator 1. Aus der letzten Formel für α_{mn} folgt sofort (l_0 beliebig, aber $< l_1$):

$$\alpha_{mn} = \sum_{p=1}^h l_p (e_{mn}(l_p) - e_{mn}(l_{p-1})) = \int_{-\infty}^{\infty} \lambda d e_{mn}(\lambda).^{15)}$$

¹³) Wie man sieht, sind $\lambda_1, \dots, \lambda_k$ Eigenwerte, und $(\overline{u_{11}}, \dots, \overline{u_{k1}}), \dots, (\overline{u_{1k}}, \dots, \overline{u_{kk}})$ dazugehörige Eigenvektoren (vgl. II.).

¹⁴) Was bei überall sinnvollen Operatoren A, B unter AB zu verstehen ist, ist klar.

¹⁵) Stieltjessches Integral!

Hieraus folgt sofort (wenn, ebenso wie in IV., $(f, g) = \sum_{n=1}^k x_n \bar{y}_n$ ist, falls $f = (x_1, \dots, x_k)$, $g = (y_1, \dots, y_k)$ ist)

$$(\beta) \quad (f, Rg) = \int_{-\infty}^{\infty} \lambda d(f, E(\lambda)g).$$

Man überlegt sich leicht, daß (α) , (β) , zusammen mit $E(\lambda) = 0$ bzw. 1 für kleine bzw. große λ , und damit, daß $E(\lambda)$ in λ nach rechts stetig ist, die Operatoren $E(\lambda)$ vollkommen festlegt. Da aber aus den $E(\lambda)$ (bzw. den $e_{mn}(\lambda)$) die Diagonalform von R leicht erzeugt werden kann, ist dies die eindeutige Formulierung des Eigenwertproblems.

Kehren wir nun zum abstrakten Hilbertschen Raume $\S$ zurück! Die Bedingungen (α) , (β) dürfen da wörtlich übertragen werden, wenn wir nur die $E(\lambda)$ genau erklären: sie sollen überall sinnvolle H. O. sein¹⁶⁾. Die Zusatzforderungen betreffs der λ -Abhängigkeit der $E(\lambda)$ formulieren wir so:

$$(\gamma) \quad \text{für } \lambda \rightarrow -\infty \quad E(\lambda) \rightarrow 0, \quad \text{für } \lambda \rightarrow +\infty \quad E(\lambda) \rightarrow 1, \\ \text{für } \lambda \geq \lambda_0, \lambda \rightarrow \lambda_0 \quad E(\lambda) \rightarrow E(\lambda_0).^{17)}$$

(Daß gerade dies die richtige Generalisierung ist, wird der Erfolg zeigen.)

Eine (α) , (γ) genügende Schar $E(\lambda)$ heiße eine Zerlegung der Einheit (kurz: Z. d. E.); wenn R zu ihr in der Beziehung (β) steht, ist es der dazugehörige Operator. Freilich entstehen Konvergenzfragen, wir werden aber sehen (vgl. Satz 36), daß diese am besten behoben werden, wenn wir (β) so verschärfen: Rg habe dann und nur dann Sinn, wenn

$$(\beta') \quad \int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)g|^2$$

endlich ist — dann sind (wie wir zeigen werden) alle Integrale von (β) absolut konvergent — und es möge (β) gelten. (Aus (α) folgt, daß $|E(\lambda)f|^2$ als Funktion von λ nie abnimmt, vgl. Satz 17, 18, der Integrand λ^2 von (β') ist ≥ 0 : also ist das Integral (β') entweder absolut konvergent, oder eigentlich divergent, und zwar $+\infty$.) Man sieht leicht ein: eine Z. d. E. (also eine die (α) , (γ) genügt) bestimmt den zugehörigen Operator (wenn es einen gibt) eindeutig: denn der Definitionsbereich ist gegeben (durch (β')), und alle (f, Rg) (durch (β)), also auch Rg .

¹⁶⁾ Im k -dimensionalen Raume hat $E(\lambda)^2 = E(\lambda)$, wie man leicht verifiziert, die Konsequenz $|E(\lambda)f| \leq |f|$; es ist also zu erwarten, daß es in $\S$ ebenso sein wird, also daß $E(\lambda)$ als beschränkt und somit überall sinnvoll anzusetzen ist (vgl. auch Kap. III, Satz 14, 15, 16).

¹⁷⁾ Bei (überall sinnvollen) Operatoren A_n, B verstehen wir unter $A_n \rightarrow B$, daß für alle f $A_n f \rightarrow Bf$ gilt. 1 ist durch $1f = f$ definiert.

Das Eigenwertproblem besteht also darin, zu entscheiden: Gibt es zu jeder Z. d. E. einen zugehörigen H. O.? Ist dieser maximal? Müssen zu verschiedenen Z. d. E. verschiedene Operatoren gehören? Welche H. O. gehören zu Z. d. E.?¹⁸⁾ Wir werden sehen, daß die drei ersten Fragen bejahend zu beantworten sind (vgl. Satz 36). Bei der vierten müssen wir daher weiter fragen: Die den Z. d. E. zugeordneten H. O. sind eine Teilklasse der Klasse aller max. H. O. Welche ist nun diese Teilklasse?

VII. Wir werden sehen (Satz 36, 38), daß sie nicht alle max. H. O. umfaßt. Es gibt eine Ausnahme¹⁹⁾ $\bar{R}$, und alle anderen Ausnahmen lassen sich durch gewisse, hier nicht zu erörternde, triviale Operationen aus ihr und den H. O. mit Z. d. E. aufbauen (vgl. Satz 38). Bei $\bar{R}$ werden wir sehen, daß seine Eigenschaften in der Tat jede vernünftige Formulierung eines Eigenwertproblems ausschließen, daß es aber keineswegs „pathologisch“ ist: es ist leicht zu beschreiben (vgl. Anm. ¹⁹⁾).

Immerhin müssen solche max. H. O., die beschränkt, oder definit, oder reell²⁰⁾ sind, eine Z. d. E. besitzen, gewisse einfache Eigenschaften sind somit mit dem Fehlen der Z. d. E. unverträglich (vgl. Satz 40). Das Auftreten dieses $\bar{R}$ können wir auch so deuten, daß der Hermitesche Charakter bei unendlichen Operatoren das Auftreten der Ausnahmefälle der Elementarteilertheorie nicht mehr verhindert (im Gegensatz zum Falle endlich vieler Dimensionen; oder im Unendlichvioldimensionalen, zu den reellen, oder beschränkten, oder definiten Operatoren), aber daß nur ein einziger Ausnahme-Typus vorkommen kann.

Es genügt wohl vorläufig soviel über die max. H. O. zu sagen. (Insbesondere soll nicht weiter erörtert werden, wie und in welchem Umfange die Kenntnis der Z. d. E. — d. h. der „Hilbertschen Spektralform“ — die Lösung des Eigenwertproblems, so wie es in I., II. formuliert wurde, ermöglicht. Es sei diesbezüglich auf die in Anm. ⁴⁾, ⁵⁾ zitierten Untersuchungen

¹⁸⁾ Hilbert bewies: Für beschränkte H. O. gibt es stets eine und nur eine Z. d. E., und zwar kommen die und nur die Z. d. E. vor, für die $E(\lambda)$ für hinreichend kleine wie für hinreichend große λ konstant (also = 0 bzw. 1) wird. (β') fällt dann fort, da es stets erfüllt ist. (Zitat: Anm. ⁷⁾.) — Bei unserer Formulierung ist offenbar die Abtrennung von Punkt- und Streckenspektrum noch nicht vollzogen (in (γ) wird $\lambda \geq \lambda_0$ verlangt!), sie ist für unsere Zwecke unerheblich.

¹⁹⁾ Der unter anderem der Anhang IV dieser Arbeit gewidmet wird.

²⁰⁾ Sei $\varphi_1, \varphi_2, \dots$ ein vollständiges Orthogonalsystem in $\S$ (vgl. Def. 3), dann kann jedes f eindeutig auf die Form $\sum_{n=1}^{\infty} x_n \varphi_n$ gebracht werden, es heiße reell, wenn alle x_n es sind. $\Re f$ entsteht aus f , wenn alle x_n durch $\Re x_n$ (Realteil!) ersetzt werden. R ist reell, wenn bei sinnvollem Rf auch $R \Re f$ sinnvoll ist, und zwar $= \Re Rf$. (In den Begriff der „Realität“ geht also noch willkürlich ein gewisses vollständiges Orthogonalsystem ein.)

verwiesen.) Dagegen ist noch der Fall solcher H. O. zu erörtern, deren Maximalität nicht gesichert (oder nicht vorhanden) ist. Die H. O. sind ja meistens so gegeben: z. B. als selbstadjungierte Differentialoperatoren, definiert für alle n -mal differentiiierbaren Funktionen (n ist der Grad!), die gewissen Randbedingungen genügen. Wir haben in V. erwähnt, daß jeder H. O. zu einem max. H. O. fortsetzbar ist, es ist aber nicht gesagt, daß dies nur auf eine einzige Weise geht.

Wir werden auch diesbezüglich einfache Resultate erzielen, so u. a.: wenn ein H. O. nicht selbst max. oder nicht nach einer gewissen trivialen Forts. max. ist, so kann er auf Kontinuum viele verschiedene Arten max. gemacht werden (vgl. die Schlußbemerkung von Kap. VIII). Beschränkte H. O. sind stets im ersten Falle.

Es werde noch an einem einfachen Beispiel klargemacht, wie diese mehrdeutige Fortsetzbarkeit zustande kommt. Sei Ω das Intervall $0, 1$, R der Operator $i \frac{d}{dx} \dots$, definiert für alle differentiiierbaren Funktionen. Da

$$\begin{aligned}(f, Rg) - (Rf, g) &= \int_0^1 i \{f(x) \bar{g}'(x) + f'(x) \bar{g}(x)\} dx \\ &= i \{f(1) \bar{g}(1) - f(0) \bar{g}(0)\}\end{aligned}$$

ist, ist R nicht Hermitesch, wohl aber R_0 und R_ϑ (ϑ eine Zahl vom Absolutwerte 1), die aus R durch die folgende Einschränkung des Definitionsbereiches entstehen:

$$f(0) = f(1) = 0 \quad \text{bzw.} \quad f(1) = \vartheta \cdot f(0).$$

R_ϑ hat max. Fortsetzungen, z. B. S_ϑ . Alle S_ϑ sind auch max. Forts. von R_0 . Und sie sind alle verschieden, denn wäre $S_{\vartheta_1} = S_{\vartheta_2} = S$, so wäre S für ein f und ein g mit $f(0) = g(0) = 1$, $f(1) = \vartheta_1$, $g(1) = \vartheta_2$ sinnvoll, und daher

$$(f, Sg) - (Sf, g) = \vartheta_1 \bar{\vartheta}_2 - 1 \neq 0,$$

also S kein H. O. Die Mehrdeutigkeit der max. Forts. rührt hier also von der Mehrzahl der möglichen Randbedingungen her. Wir werden sehen, daß noch im allgemeinsten Falle die Verhältnisse eine gewisse Analogie hiermit aufweisen (vgl. Satz 26, 27)²¹⁾.

VIII. Der Zusammenhang der H. O. von § mit den Matrizen ist bekanntlich dieser: sei R ein H. O. und $\varphi_1, \varphi_2, \dots$ ein vollständiges Orthogonalsystem (vgl. Def. 3), dann wird R die Matrix $\{\alpha_{mn}\}$, $\alpha_{mn} = (\varphi_m, R \varphi_n)$, zugeordnet. Diese Zuordnung ist bei beschränkten H. O. ganz klar, bei

²¹⁾ Bei singulären Integraloperatoren fand schon Carleman Zeichen dieser Mehrdeutigkeit, jedoch konnte er in diesem Falle das Eigenwertproblem nicht ganz lösen (es fehlte das Analogon der entscheidenden Bedingung (α) aus IV.). Vgl. Anm. 4).

unbeschränkten ist sie mit gewissen Komplikationen verknüpft. Sie wird uns im Anhang III beschäftigen. Die unitäre Transformationstheorie der unbeschränkten Matrizen führt in eine besonders eigentümliche Pathologie, dieselbe soll aber an anderer Stelle untersucht werden²²⁾.

Auch über das Verhältnis dieser Resultate zur neuen Quantentheorie (der sogenannten „Transformationstheorie“) wollen wir hier nicht sprechen, jedenfalls zeigen sie, daß die allgemeine Theorie der H. O. keineswegs überall dasjenige Verhalten zeigt, welches in der „Transformationstheorie“ (auf Grund der Analogie mit beschränkten oder gar endlichvioldimensionalen Operatoren) im allgemeinen vermutet und vorweggenommen wurde.

IX. Es bleibt noch übrig, einiges über die Methode und Anordnung dieser Arbeit zu sagen.

Der wesentliche Kunstgriff ist das Untersuchen eines gewissen Operators U , der symbolisch mit $U = \frac{R+i1}{R-i1}$ bezeichnet werden könnte, an Stelle von R . Wenn R ein H. O. ist, so ist nämlich U längentreu (d. h. $|Uf| = |f|$), also stetig, und somit leichter zu behandeln (vgl. Satz 22, 23, 24). Diese „Cayley-Transformation“ $U = \frac{R+i1}{R-i1}$ ist der Kern unserer Methode.

Alle anderen Methoden versagen: Die eleganten Verfahren von Hellinger und von F. Riesz, weil in ihnen der Operator R iteriert werden muß, was nur bei überall sinnvollen (also beschränkten) H. O. ohne weiteres geht, bei beliebigen H. O. aber zunächst fraglich ist. Alle Maximum-Minimum-Methoden, sowie die Methode von E. Schmidt sind von vornherein auf Fälle ohne Streckenspektrum beschränkt. Die ursprüngliche Hilbertsche Methode (Approximation durch endlichvioldimensionale Abschnittoperatoren) ist es allein, die gewisse Resultate zu erzielen erlaubt: aber nur unter sehr großen Schwierigkeiten, und nur einen Bruchteil dessen, was wir erreichen können²³⁾.

²²⁾ Vgl. eine demnächst im Journal f. Math. erscheinende Arbeit des Verfassers: „Zur Theorie der unbeschränkten Matrizen“.

²³⁾ Der Verfasser konnte mit dieser Methode den Fall reeller H. O. (vgl. Anm. ²⁰⁾) erledigen, wo der in VII. angedeutete Ausnahmefall noch nicht in Erscheinung trat. Dabei konnte aber der Fortsetzungsprozeß der H. O. nicht so übersehen werden, wie es etwa im Kap. VIII der Fall sein wird, und die Methode war derart unkonstruktiv, daß dabei z. B. der Wohlordnungssatz herangezogen werden mußte (vgl. die Ankündigung im Jahresber. d. D. Math.-Ver. 37, 1–4 (1928), S. 11–15), auch war den nicht-reellen H. O. so nicht beizukommen. — Der Verfasser möchte es nicht versäumen, Herrn E. Schmidt an dieser Stelle seinen Dank für das Interesse auszusprechen, das er diesen Untersuchungen (insbesondere der Übertragung der Resultate auf nicht-reelle H. O.) zugewandt hat, und darauf hinweisen, daß er wiederholt wesentliche Anregungen aus Gesprächen mit ihm empfing. Insbesondere stammt der Begriff der „Hypermaximalität“ von H. O. (vgl. Def. 9), sowie die Erkenntnis, daß diese Eigenschaft für die Existenz der Z. d. E. notwendig und hinreichend ist, von E. Schmidt.

Die Anordnung des Gegenstandes ist die folgende: Die Kap. I bis IV sind vorbereitender Natur: I. Struktur des Raumes §, II. Allgemeines über Operatoren in §, III. Lineare Mannigfaltigkeiten und ihre Projektionsoperatoren, IV. Reduzibilität von Operatoren. In den Kap. V bis VIII wird die Hauptarbeit der Lösung des Eigenwertproblems geleistet: V. Cayleysche Transformation, VI. Erweiterungselemente, VII. Defektindizes, VIII. Der Fortsetzungsprozeß. In Kap. IX bis X werden im Anschluß daran die genauen Normalformen für die max. H. O. hergestellt: IX. Eigenwertdarstellung, X. Operatoren ohne Eigenwertdarstellung. Kap. XI bis XII enthalten schließlich einige allgemeine Sätze über unbeschränkte Operatoren: XI. Halbbeschränktheit, XII. Pathologie der Unbeschränktheit²⁴). Außerdem sind noch einige Anhänge da: In I. wird bewiesen, daß die Funktionenräume von I. (Einleitung) wirklich den Bedingungen A bis E von Kap. I genügen (also daß unsere Theorie auf sie anwendbar ist)²⁵). In Anhang II beweisen wir einen Satz über die Lösung des Eigenwertproblems unitärer Operatoren (allgemeiner dessen von „normalen“) — dies ist eine Frage über beschränkte Operatoren, die mit der F. Rieszschen Methode lösbar ist, aber im Rahmen der Hilbertschen Theorie bisher nie erledigt wurde²⁶). Wir brauchen dies für Kap. VIII, da es aber methodisch in den Kreis der Hilbertschen Theorie gehört, ziehen wir vor, es so getrennt zu behandeln. In Anhang III gehen wir auf den Zusammenhang von Operatoren und Matrizen näher ein (vgl. VIII. Einleitung, sowie Anm.²³); in IV wird der Operator $\bar{R}$ (vgl. VII. Einleitung) genauer untersucht.

Im übrigen ist diese Arbeit ein logisch geschlossenes Ganzes, das ohne alle Vorkenntnisse aus der Literatur der „unendlich vielen Variablen“ gelesen werden kann. (Dementsprechend ist auch der Inhalt der Kap. I, III und des Anhangs I — abgesehen von der Anordnung — nicht neu.)

I. Der abstrakte Hilbertsche Raum.

Wir charakterisieren den abstrakten (komplexen) Hilbertschen Raum durch fünf Eigenschaften A bis E. Es erweist sich als zweckmäßig, an einige von ihnen einige einfache Sätze anzuschließen, die bereits aus den betreffenden allein folgen.

²⁴) Dieselbe wird in der in Anm.²²) angekündigten Arbeit des Verfassers noch bedeutend verschärft werden.

²⁵) Dies ergibt, zusammen mit einigen Überlegungen von Kap. I, natürlich einen Beweis des Fischer-F. Rieszschen Satzes.

²⁶) Es kommt im wesentlichen darauf heraus, daß zwei vertauschbare beschränkte H. O. eine gemeinsame Eigenwertform (oder Spektralform) nach Hilbert zulassen. Vgl. Hellinger-Toeplitz, Encyklopädie-Artikel II, C. 13, S. 1562 u. f., wo die Frage für vollstetige H. O. gelöst wird.

A. $\mathfrak{S}$ ist ein linearer Raum.

D. h.: es gibt in $\mathfrak{S}$ eine Addition $f + g$, eine Subtraktion $f - g$, und eine (skalare) Multiplikation af (f, g von $\mathfrak{S}$, a eine komplexe Zahl), sowie ein Element 0 ; und für diese gelten die Rechenregeln der gewöhnlichen Vektoralgebra.

Definition 1. Eine Teilmenge $\mathfrak{M}$ von $\mathfrak{S}$ heißt eine lineare Mannigfaltigkeit (kurz: lin. M.), wenn sie mit $f_1, \dots, f_n$ auch $a_1 f_1 + \dots + a_n f_n$ enthält. Wenn $\mathfrak{M}$ irgendeine Teilmenge von $\mathfrak{S}$ ist, so ist die Menge aller $a_1 f_1 + \dots + a_n f_n$ ($f_1, \dots, f_n$ von $\mathfrak{M}$, $a_1, \dots, a_n$ Konstante) die kleinste $\mathfrak{M}$ enthaltende lin. M., sie heiße die von $\mathfrak{M}$ aufgespannte lin. M.

Definition 2. n Elemente $f_1, \dots, f_n$ heißen linear unabhängig (kurz: lin. unabh.), wenn aus $a_1 f_1 + \dots + a_n f_n = 0$ $a_1 = \dots = a_n = 0$ folgt. n lin. M. $\mathfrak{M}_1, \dots, \mathfrak{M}_n$ heißen lin. unabh., wenn aus $f_1 + \dots + f_n = 0$ (f_1 von $\mathfrak{M}_1, \dots, f_n$ von $\mathfrak{M}_n$) $f_1 = \dots = f_n = 0$ folgt.

B. *Es gibt in $\mathfrak{S}$ ein, zu dem der Vektorrechnung analoges, inneres Produkt, das eine Metrik erzeugt.*

D. h.: es gibt eine Funktion (f, g) (f, g von $\mathfrak{S}$, (f, g) eine komplexe Zahl) mit den folgenden Eigenschaften:

1. $(af, g) = a(f, g)$, 2. $(f_1 + f_2, g) = (f_1, g) + (f_2, g)$, 3. $(f, g) = \overline{(g, f)}$,
4. $(f, f) \geq 0$ und nur für $f = 0$ ist es $= 0$.

(Aus 1., 2. folgt nach 3.: 1'. $(f, ag) = \bar{a}(f, g)$, 2. $(f, g_1 + g_2) = (f, g_1) + (f, g_2)$.) Durch

$$|f| = \sqrt{(f, f)}$$

wird ein „Betrag“ definiert, durch $|f - g|$ die Metrik (vgl. Satz 2 und Anm.²⁷⁾).

Satz 1. *Es gilt stets*

$$|(f, g)| \leq |f| \cdot |g|.$$

Beweis. Erstens ist (nach 2., 2', 3. und 4.)

$$\begin{aligned} \Re(f, g) &= \left(\frac{f+g}{2}, \frac{f+g}{2}\right) - \left(\frac{f-g}{2}, \frac{f-g}{2}\right) \leq \left(\frac{f+g}{2}, \frac{f+g}{2}\right) + \left(\frac{f-g}{2}, \frac{f-g}{2}\right) \\ &= \frac{1}{2}((f, f) + (g, g)) = \frac{|f|^2 + |g|^2}{2}. \end{aligned}$$

Wir ersetzen f, g durch $af, \frac{1}{a}g$ ($a > 0$): die linke Seite bleibt invariant, die rechte hat das Minimum $|f| \cdot |g|$. Wir ersetzen f durch $\vartheta \cdot f$ ($|\vartheta| = 1$), die rechte Seite bleibt invariant, die linke hat das Maximum $|(f, g)|$. Also

$$|(f, g)| \leq |f| \cdot |g|.$$

Satz 2. *Es gilt stets:*

$$|af| = |a| \cdot |f|, \quad |f+g| \leq |f| + |g|.$$

Beweis. Das erste folgt sofort aus 1., 1', das zweite aus Satz 1:

$$(f+g, f+g) = (f, f) + (g, g) + 2\Re(f, g),$$

$$|f+g|^2 = |f|^2 + |g|^2 + 2\Re(f, g) \leq |f|^2 + |g|^2 + 2|f| \cdot |g|,$$

$$|f+g| \leq |f| + |g|.$$

Definition 3. Zwei f, g sind orthogonal (kurz: orth.), wenn $(f, g) = 0$ ist; zwei lin. M. sind orth., wenn jedes Element der einen zu jedem der anderen orth. ist. Eine Menge $\mathfrak{M}$ heißt normiert orthogonal (kurz: norm. orth.), wenn

$$(f, g) = \begin{cases} 1, & \text{für } f = g \\ 0, & \text{für } f \neq g \end{cases} \quad f, g \text{ von } \mathfrak{M}$$

gilt. Sie ist vollständig (kurz: vollst. norm. orth.), wenn sie nicht echter Teil einer anderen norm. orth. Menge ist, — was offenbar besagt, daß $(f, g) = 0$ (f fest, für alle g von $\mathfrak{M}$) $|f| = 1$ ausschließt, also auch $|f| > 0$, also $f = 0$ zur Folge hat. (Bekanntlich hat bei mehreren Elementen oder lin. M. die paarweise Orth. die lin. Unabh. zur Folge.)

Definition 4. Nach Satz 2 ist $|f-g|$ eine Entfernung in $\mathfrak{S}^{27}$, also $\mathfrak{S}$ topologisiert. Eine abgeschlossene (kurz: abg.) lin. M. ist also eine, die ihre Häufungspunkte enthält. Wenn wir zur von $\mathfrak{M}$ aufgespannten lin. M. alle ihre Häufungspunkte hinzufügen, erhalten wir die kleinste $\mathfrak{M}$ enthaltende abg. lin. M.: die von $\mathfrak{M}$ aufgespannte abg. lin. M.

C. In der Metrik $|f-g|$ ist $\mathfrak{S}$ separabel. D. h.: eine gewisse abzählbare Menge ist in $\mathfrak{S}$ überall dicht.

D. $\mathfrak{S}$ besitzt beliebig (endlich!) viele lin. unabh. Elemente²⁸).

²⁷) Von einer „Entfernung“ $\overline{f, g}$ verlangt man bekanntlich, daß sie den folgenden Axiomen genüge:

1. $\overline{f, g} \geq 0$, und nur für $f = g$ ist es $= 0$;

2. $\overline{f, g} = \overline{g, f}$;

3. $\overline{f, h} \leq \overline{f, g} + \overline{g, h}$.

In linearen Räumen kommt noch

4. $\overline{f+h, g+h} = \overline{f, g}, \quad \overline{af, ag} = |a| \cdot \overline{f, g}$

hinzu. — Alle diese Eigenschaften besitzt $\overline{f, g} = |f-g|$ nach Satz 2. Man beachte, daß (f, g) wegen seiner Bilinearität (1., 1' in B) und Satz 1 stetig ist.

²⁸) Hierdurch wird die Möglichkeit, daß $\mathfrak{S}$ ein endlichvieldegmentaler Euklidischer Raum sei, ausgeschlossen. Die Bedingungen A, B sind übrigens in Anlehnung an ein (für endlich viele Dimensionen aufgestelltes) Axiomensystem des Vektorraumes von H. Weyl (Raum, Zeit, Materie, 5. Aufl., Berlin 1923, §§ 2—4) gebildet.

E. $\mathfrak{H}$ ist vollständig. D. h.: wenn eine Folge $f_1, f_2, \dots$ in $\mathfrak{H}$ der Cauchyschen Konvergenzbedingung genügt (zu jedem $\varepsilon > 0$ gibt es ein $N = N(\varepsilon)$, so daß aus $m, n \geq N$ $|f_m - f_n| \leq \varepsilon$ folgt), so ist sie konvergent (es existiert ein f aus $\mathfrak{H}$, so daß es zu jedem $\varepsilon > 0$ ein $N = N(\varepsilon)$ gibt, so daß aus $n \geq N$ $|f_n - f| \leq \varepsilon$ folgt).

Im Anhang I werden wir zeigen, daß alle in der Einleitung I. und II. genannten Funktionenräume die Eigenschaften A bis E haben, also Spezialfälle des abstrakten Hilbertschen Raumes $\mathfrak{H}$ sind. Hier aber gehen wir daran zu beweisen, daß $\mathfrak{H}$ in der Tat durch die Eigenschaften A bis E vollständig charakterisiert ist: daß $\mathfrak{H}$ mit dem gewöhnlichen Hilbertschen Raume isomorph ist. Dazu ist es notwendig, einige allgemein bekannte Sätze über norm. orth. und vollst. norm. orth. Systeme zu beweisen.

Satz 3. Jede norm. orth. Menge $\mathfrak{M}$ ist höchstens abzählbar, jede vollst. norm. orth. Menge $\mathfrak{M}$ ist genau abzählbar.

Bemerkung. Wir schreiben darum von nun an alle norm. Orth. und vollst. norm. orth. Mengen als Folgen $\varphi_1, \varphi_2, \dots$ (die im ersteren Falle eventuell abbrechen).

Beweis. $\mathfrak{M}$ sei norm. orth., für zwei beliebige f, g von $\mathfrak{M}$, $f \neq g$ gilt

$$|f - g|^2 = |f|^2 + |g|^2 - 2 \Re(f, g) = 2, \quad |f - g| = \sqrt{2};$$

wegen der Separabilität von $\mathfrak{H}$ (Bedingung C) ist also $\mathfrak{M}$ höchstens abzählbar.

Wir müssen noch zeigen, daß ein endliches norm. orth. System $\mathfrak{M}$ nicht vollst. ist. Seien die Elemente von $\mathfrak{M}$ $\varphi_1, \dots, \varphi_n$. Unter den $a_1 \varphi_1 + \dots + a_n \varphi_n$ gibt es keine $n + 1$ lin. unabh., also muß $\mathfrak{H}$ ein Element f haben, das von ihnen allen verschieden ist (Bedingung D). Dann ist

$$f^* = f - a_1 \varphi_1 - \dots - a_n \varphi_n$$

jedenfalls $\neq 0$, und zu allen $\varphi_1, \dots, \varphi_n$ orth., wenn wir $a_k = (f, \varphi_k)$ ($k = 1, \dots, n$) setzen. Nach der Schlußbemerkung von Def. 3 ist also $\mathfrak{M}$ unvollst.

Satz 4. $\varphi_1, \varphi_2, \dots$ sei ein norm. orth. System. Dann ist die Reihe

$$\sum_{n=1}^{\infty} (f, \varphi_n) \cdot \overline{(g, \varphi_n)}$$

für alle f, g von $\mathfrak{H}$ absolut konvergent, insbesondere ist für $f = g$

$$\sum_{n=1}^{\infty} |(f, \varphi_n)|^2 \leq |f|^2.$$

Beweis. Für $a_n = (f, \varphi_n)$ gilt bekanntlich:

$$\begin{aligned} |f - \sum_{n=1}^N a_n \varphi_n|^2 &= |f|^2 - 2 \Re \sum_{n=1}^N (f, a_n \varphi_n) + \sum_{m,n=1}^N (a_m \varphi_m, a_n \varphi_n) \\ &= |f|^2 - 2 \Re \sum_{n=1}^N \overline{a_n} (f, \varphi_n) + \sum_{m,n=1}^N a_m \overline{a_n} (\varphi_m, \varphi_n) \\ &= |f|^2 - 2 \sum_{n=1}^N |a_n|^2 + \sum_{n=1}^N |a_n|^2 = |f|^2 - \sum_{n=1}^N |a_n|^2. \end{aligned}$$

Da die linke Seite stets ≥ 0 ist, haben wir:

$$\sum_{n=1}^N |a_n|^2 \leq |f|^2, \quad \sum_{n=1}^N |(f, \varphi_n)|^2 \leq |f|^2.$$

Somit ist $\sum_{n=1}^{\infty} |(f, \varphi_n)|^2$ konvergent (absolut!), und zwar $\leq |f|^2$. Und weil

$$|(f, \varphi_n) \cdot \overline{(g, \varphi_n)}| \leq \frac{1}{2} |(f, \varphi_n)|^2 + \frac{1}{2} |(g, \varphi_n)|^2$$

ist, konvergiert auch $\sum_{n=1}^{\infty} (f, \varphi_n) \cdot \overline{(g, \varphi_n)}$ absolut. Damit ist alles bewiesen.

Satz 5. $\varphi_1, \varphi_2, \dots$ sei ein norm. orth. System. Die Reihe $\sum_{n=1}^{\infty} a_n \varphi_n$ konvergiert dann und nur dann, wenn $\sum_{n=1}^{\infty} |a_n|^2$ konvergiert.

Beweis. Nach E genügt es zu zeigen, daß die Cauchysche Konvergenzbedingung für beide Reihen gleich lautet: Dies ist aber wegen ($m \geq n$)

$$\begin{aligned} \left| \sum_{p=1}^m a_p \varphi_p - \sum_{p=1}^n a_p \varphi_p \right|^2 &= \left| \sum_{p=n+1}^m a_p \varphi_p \right|^2 = \sum_{p,q=n+1}^m (a_p \varphi_p, a_q \varphi_q) \\ &= \sum_{p,q=n+1}^m a_p \overline{a_q} (\varphi_p, \varphi_q) = \sum_{p=n+1}^m |a_p|^2 = \sum_{p=1}^m |a_p|^2 - \sum_{p=1}^n |a_p|^2 \end{aligned}$$

in der Tat der Fall.

Zusatz. Für $f = \sum_{n=1}^{\infty} a_n \varphi_n$ gilt (im Falle der Konvergenz)

$$(f, \varphi_n) = a_n.$$

Beweis. Für $N \geq n$ ist

$$\left(\sum_{p=1}^N a_p \varphi_p, \varphi_n \right) = \sum_{p=1}^N a_p (\varphi_p, \varphi_n) = a_n,$$

und wegen der Stetigkeit des inneren Produktes dürfen wir $N \rightarrow \infty$ lassen, was $(f, \varphi_n) = a_n$ ergibt.

Satz 6. $\varphi_1, \varphi_2, \dots$ sei ein norm. orth. System. Für jedes f ist die Reihe

$$f' = \sum_{n=1}^{\infty} a_n \varphi_n, \quad a_n = (f, \varphi_n) \quad (n = 1, 2, \dots)$$

konvergent, und zwar ist $f - f'$ zu allen $\varphi_1, \varphi_2, \dots$ orth.

Beweis. Die Reihe konvergiert nach Satz 4 und 5, und nach dem Zusatz zu Satz 5 ist $(f', \varphi_n) = a_n = (f, \varphi_n)$, $(f - f', \varphi_n) = 0$; also ist alles bewiesen.

Satz 7. *Dazu, daß ein norm. orth. System $\varphi_1, \varphi_2, \dots$ vollst. sei, ist eine jede der drei folgenden Bedingungen notwendig und hinreichend:*

$\alpha)$ $\varphi_1, \varphi_2, \dots$ spannen die abg. lin. M. $\mathfrak{H}$ auf.

$\beta)$ Es ist für alle f

$$f = \sum_{n=1}^{\infty} a_n \varphi_n, \quad a_n = (f, \varphi_n) \quad (n = 1, 2, \dots).$$

$\gamma)$ Es ist für alle f, g

$$(f, g) = \sum_{n=1}^{\infty} (f, \varphi_n) \cdot \overline{(g, \varphi_n)}.$$

Beweis. Erstens folgt aus der Vollst. $\beta)$: denn $f - f'$ aus Satz 6 muß verschwinden. Zweitens folgt aus $\beta)$ $\alpha)$, denn f ist Häufungspunkt der $\sum_{n=1}^N a_n \varphi_n$, d. h. der von $\varphi_1, \varphi_2, \dots$ aufgespannten lin. M. — also Element der abg. lin. M. Drittens folgt aus $\beta)$ $\gamma)$, denn es ist

$$\sum_{n=1}^N a_n \varphi_n \rightarrow f, \quad \left| \sum_{n=1}^N a_n \varphi_n \right|^2 \rightarrow |f|^2 \quad (\text{für } N \rightarrow \infty),$$

$$\left| \sum_{n=1}^N a_n \varphi_n \right|^2 = \sum_{m,n=1}^N a_m \overline{a_n} (\varphi_m, \varphi_n) = \sum_{n=1}^N |a_n|^2,$$

also

$$\sum_{n=1}^{\infty} |a_n|^2 = \sum_{n=1}^{\infty} |(f, \varphi_n)|^2 = |f|^2.$$

Wenn wir hier f durch $\frac{f+g}{2}$ und $\frac{f-g}{2}$ ersetzen und subtrahieren, gewinnen wir den Realteil von $\gamma)$, wenn wir f, g durch if, g ersetzen, ihren Imaginärteil.

Viertens folgt aus $\alpha)$ die Vollst., denn sonst gäbe es ein zu allen $\varphi_1, \varphi_2, \dots$ orth. $f \neq 0$. Also ist die ganze von $\varphi_1, \varphi_2, \dots$ aufgespannte abg. lin. M., d. h. $\mathfrak{H}$, zu f orth. — also auch f : $|f|^2 = (f, f) = 0$, $f = 0$, entgegen der Annahme. Fünftens folgt auch aus $\gamma)$ die Vollst.: denn für dasselbe f muß nach $\gamma)$ $(f = g)$

$$|f|^2 = (f, f) = \sum_{n=1}^{\infty} 0^2 = 0, \quad f = 0$$

sein.

Wir haben also das folgende logische Schema:

$$\text{Vollständigkeit} \rightarrow \beta \begin{matrix} \nearrow \alpha \\ \searrow \gamma \end{matrix} \rightarrow \text{Vollständigkeit.}$$

Folglich sind diese vier Aussagen in der Tat gleichbedeutend.

Satz 8. Zu jeder Folge $f_1, f_2, \dots$ beliebiger Elemente von $\mathfrak{H}$ gibt es ein norm. orth. System $\varphi_1, \varphi_2, \dots$, welches dieselbe lin. M. (also auch dieselbe abg. lin. M.) aufspannt (eine jede der beiden Folgen kann bereits im Endlichen abbrechen).

Bemerkung. Wir wählen $f_1, f_2, \dots$ überall dicht (Bedingung C!), dann spannt es natürlich die abg. lin. M. $\mathfrak{H}$ auf. Ebenso das nach Satz 8 konstruierte norm. orth. System. $\varphi_1, \varphi_2, \dots$, nach Satz 7 α) ist es somit vollst. Damit ist die Existenz von vollst. norm. orth. Systemen erwiesen.

Beweis. Erstens können wir $f_1, f_2, \dots$ durch eine Teilfolge von lauter lin. unabh. Elementen ersetzen, die dieselbe lin. M. aufspannt. In der Tat sei g_1 das erste $f_n \neq 0$, g_2 das erste $f_n \neq a_1 g_1$ (a_1 beliebig), g_3 das erste $f_n \neq a_1 g_1 + a_2 g_2$ (a_1, a_2 beliebig) usw. (wenn es kein $f_n \neq a_1 g_1 + \dots + a_p g_p$ gibt, brechen wir mit g_p ab). $g_1, g_2, \dots$ leistet offenbar das zunächst Geforderte.

Auf $g_1, g_2, \dots$ wenden wir nunmehr das E. Schmidtsche „Orthogonalisationsverfahren“ an:

$$\begin{aligned} \gamma_1 &= g_1, & \varphi_1 &= \frac{1}{|\gamma_1|} \cdot \gamma_1. \\ \gamma_2 &= g_2 - (g_2, \varphi_1) \cdot \varphi_1, & \varphi_2 &= \frac{1}{|\gamma_2|} \cdot \gamma_2, \\ \gamma_3 &= g_3 - (g_3, \varphi_1) \cdot \varphi_1 - (g_3, \varphi_2) \cdot \varphi_2, & \varphi_3 &= \frac{1}{|\gamma_3|} \cdot \gamma_3, \\ &\dots & &\dots \end{aligned}$$

$\varphi_1, \varphi_2, \dots$ spannen offenbar dieselbe lin. M. auf wie $g_1, g_2, \dots$, d. h. wie $f_1, f_2, \dots$; ferner sieht man sofort, daß sie ein norm. orth. System bilden.

Auf Grund dieser Tatsachen können wir nun zeigen, daß $\mathfrak{H}$ ein-eindeutig auf den gewöhnlichen Hilbertschen Raum, d. i. die Menge aller Folgen komplexer Zahlen $(x_1, x_2, \dots)$ mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$, abgebildet werden kann, und zwar derart, daß ($\leftrightarrow$ bezeichnet die Zuordnung) $f \leftrightarrow (x_1, x_2, \dots)$ die Folge $af \leftrightarrow (ax_1, ax_2, \dots)$ hat, $f \leftrightarrow (x_1, x_2, \dots)$, $g \leftrightarrow (y_1, y_2, \dots)$ die Folge $f + g \leftrightarrow (x_1 + y_1, x_2 + y_2, \dots)$, und außerdem $(f, g) = \sum_{n=1}^{\infty} x_n \overline{y_n}$.

Sei nämlich $\varphi_1, \varphi_2, \dots$ ein (beliebiges, aber festes) vollst. norm. orth. System, dann ordnen wir f die Folge $(x_1, x_2, \dots)$ mit $x_n = (f, \varphi_n)$ ($n = 1, 2, \dots$) zu. Daß $(x_1, x_2, \dots)$ zum Hilbertschen Raume gehört, folgt aus Satz 4, daß zu jedem solchen $(x_1, x_2, \dots)$ ein f existiert, aus Satz 5 und dessen Zusatz. Die Ein-eindeutigkeit folgt aus der Vollst. (wenn stets $(f, \varphi_n) = (g, \varphi_n)$ ist, so ist $f - g$ zu allen φ_n orth., also $= 0$). Von den drei Schlußbehauptungen sind die zwei ersten klar, die dritte folgt aus Satz 7 γ).

Somit ist gezeigt, daß $\mathfrak{H}$ dem gewöhnlichen Hilbertschen Raume isomorph ist, also zwei A bis E genügende Systeme $\mathfrak{H}$ einander; d. h. A bis E sind ein „kategorisches“ Axiomensystem, sie legen alle Eigenschaften von $\mathfrak{H}$ fest.

II. Operatoren in $\mathfrak{H}$.

Ein Operator ist eine Funktion, die in einer Teilmenge von $\mathfrak{H}$ definiert ist (eventuell in ganz $\mathfrak{H}$) und Werte aus $\mathfrak{H}$ annimmt. Wir wollen zunächst einige besonders wichtige Sorten von Operatoren definieren.

Definition 5. Ein Operator R ist linear (kurz: lin.), wenn bei sinnvollen $Rf_1, \dots, Rf_n$ auch $R(a_1f_1 + \dots + a_nf_n)$ sinnvoll, und zwar $= a_1Rf_1 + \dots + a_nRf_n$ ist. Ein Operator R ist abgeschlossen (kurz: abg.), wenn für jede Folge $f_1, f_2, \dots$, für die alle Rf_n sinnvoll sind, die Konvergenzen $f_n \rightarrow f$, $Rf_n \rightarrow f^*$ die Konsequenz haben, daß Rf sinnvoll und zwar $= f^*$ ist²⁹⁾.

Definition 6. Ein Operator R ist Hermitesch (kurz: ein H.O.), wenn erstens sein Definitionsbereich die abg lin. M. $\mathfrak{H}$ aufspannt³⁰⁾, und zweitens (soweit Rf, Rg Sinn haben)

$$(f, Rg) = (Rf, g)$$

gilt. Ein Operator R ist längentreu, wenn er abg. lin. ist und (soweit Rf Sinn hat) $|f| = |Rf|$ gilt.

Wir beginnen damit, daß wir zwei Sätze über H.O. und längentreue O. beweisen. Wir erinnern noch daran, daß wir einen Operator S Fortsetzung eines Operators R nennen (kurz: Forts.), wenn überall wo Rf Sinn hat, auch Sf sinnvoll und zwar $= Rf$ ist — und eigentliche (kurz: eig.) Forts., wenn $R \neq S$ ist, d. h. S an mehr Stellen Sinn hat, wie R .

Satz 9. Sei R ein H.O. Es gibt einen lin. H.O. $\hat{R}$, der Forts. von R ist und von dem alle ebensolchen H.O. ihrerseits Forts. sind; und es gibt einen abg. lin. H.O. $\tilde{R}$, der Forts. von R (also auch von $\hat{R}$) ist, und von dem alle ebensolchen H.O. ihrerseits Forts. sind.

Beweis. Wenn wir $\hat{R}$ so definieren könnten: $\hat{R}f$ hat Sinn, wenn $f = a_1f_1 + \dots + a_nf_n$ ist, $Rf_1, \dots, Rf_n$ sinnvoll, und zwar ist es dann $= a_1Rf_1 + \dots + a_nRf_n$; und $\tilde{R}$ so: $\tilde{R}f$ hat Sinn, wenn eine Folge

²⁹⁾ Dies ist offenbar eine stetigkeits-ähnliche Eigenschaft, und zwar in jedem „kompakten“ topologischen Raume der Stetigkeit gleichwertig. Im Hilbertschen Raume bedeutet sie aber viel weniger, für H.O. ist sie im wesentlichen immer da, vgl. Satz 9.

³⁰⁾ So weit schwächen wir (mit Rücksicht auf die Anwendung auf Matrizen, vgl. Anhang III) die Bedingung γ) von Einleitung V ab: dort wurde Überalldichtheit des Definitionsbereiches verlangt.

$g_1, g_2, \dots$ existiert mit lauter sinnvollen $\hat{R}g_n$, wobei $g_n \rightarrow f$, $\hat{R}g_n \rightarrow f^*$, und zwar ist es dann $= f^*$ — dann würden diese $\hat{R}, \tilde{R}$ offenbar alles Gewünschte leisten. Es fragt sich aber, ob diese Festsetzungen widerspruchsfrei möglich sind.

Sie sind es gewiß, wenn $a_1 f_1 + \dots + a_m f_m = b_1 g_1 + \dots + b_n g_n$ ($Rf_1, \dots, Rf_m, Rg_1, \dots, Rg_n$ sinnvoll) $a_1 Rf_1 + \dots + a_m Rf_m = b_1 Rg_1 + \dots + b_n Rg_n$ nach sich zieht; und ebenso $f_n \rightarrow f$, $g_n \rightarrow f$ (alle $\hat{R}f_n, \hat{R}g_n$ sinnvoll), $\hat{R}f_n \rightarrow f^*$, $\hat{R}g \rightarrow g^*$ die Konsequenz $f^* = g^*$ hat. Es genügt offenbar zu beweisen: Aus $a_1 f_1 + \dots + a_n f_n = 0$ ($Rf_1, \dots, Rf_n$ sinnvoll), folgt $a_1 Rf_1 + \dots + a_n Rf_n = 0$, aus $f_n \rightarrow 0$ (alle $\hat{R}f_n$ sinnvoll), $\hat{R}f_n \rightarrow f^*$ folgt $f^* = 0$.

Gelte zunächst die erste Prämisse. Wenn Rg sinnvoll ist, so ist

$$(g, a_1 Rf_1 + \dots + a_n Rf_n) = \sum_{p=1}^n (g, a_p Rf_p) = \sum_{p=1}^n (Rg, a_p f_p) \\ = (Rg, a_1 f_1 + \dots + a_n f_n) = 0.$$

Alle diese g stehen also zu $a_1 Rf_1 + \dots + a_n Rf_n$ orth., also auch ihre abg. lin. M., $\S$ — also ist $a_1 Rf_1 + \dots + a_n Rf_n$ auch zu sich selbst orth., d. h. $= 0$. Betrachten wir nunmehr die zweite Prämisse. Wenn $\hat{R}g$ Sinn hat, so ist ($\hat{R}$ ist ja ein H. O.)

$$(g, f^*) = (g, \text{limes } \hat{R}f_n) = \text{limes } (g, \hat{R}f_n) = \text{limes } (\hat{R}g, f_n) \\ = (\hat{R}g, \text{limes } f_n) = 0.$$

Somit stehen alle diese g zu f^* orth., also auch ihre abg. lin. M., $\S$ — und dies hat wieder $f^* = 0$ zur Folge.

Satz 10. *Sei U ein längentreuer O. Dann ist sowohl der Definitionsbereich als auch der Wertevorrat von U je eine abg. lin. M., und U bildet sie stetig und ein-eindeutig aufeinander ab. Soweit U Sinn hat, gilt $(f, g) = (Uf, Ug)$.*

Bemerkung. Wenn Definitionsbereich wie Wertevorrat beide $= \S$ sind, so heiße U , im Sinne der üblichen Terminologie, unitär.

Beweis. Da U lin. ist, sind Definitionsbereich und Wertevorrat lin. M. Weiter ist

$$|Uf - Ug| = |U(f - g)| = |f - g|,$$

also hat $Uf = Ug$ die Folge $f = g$, d. h. die Abbildung ist ein-eindeutig, ferner folgt hieraus ihre und ihrer Inversen Stetigkeit. Dies verursacht, zusammen mit der Abg. von U , daß Definitionsbereich und Wertevorrat beide abg. Mengen sind.

Wenn wir in $|f|^2 = |Uf|^2$ f durch $\frac{f+g}{2}$ und $\frac{f-g}{2}$ ersetzen und subtrahieren, so wird $\Re(f, g) = \Re(Uf, Ug)$. Wenn wir f, g durch if, g er-

setzen, erhalten wir dasselbe für die Imaginärteile. Also ist $(f, g) = (Uf, Ug)$. Damit ist alles bewiesen. --

Einige Eigenschaften von $\tilde{R}$ (R ein H. O.) sind evident: $\tilde{\tilde{R}} = \tilde{R}$, wenn S Forts. von R ist, so ist $\tilde{S}$ Forts. von $\tilde{R}$. Wir definieren nun die Maximalität von H. O.

Definition 7. Ein H. O. R heiße maxima! (kurz: max.), wenn es keine eig. Forts. von R gibt (die auch H. O. ist!).

Da $\tilde{R}$ Forts. von R ist, muß für max. R $\tilde{R} = R$ sein. Ferner sei R ein H. O., S max. H. O., S Forts. von R ; dann ist $S = \tilde{S}$ auch Forts. von $\tilde{R}$. D. h. jeder max. H. O. ist abg. lin., und ein H. O. R hat genau dieselben max. Forts., wie der abg. lin. H. O. $\tilde{R}$. Wir werden in den späteren Kapiteln die Forts. beliebiger H. O. zu max. untersuchen (vgl. auch die Einleitung VI., VII.), das soeben Gesagte zeigt, daß wir uns dabei durchweg auf abg. lin. Operatoren beschränken dürfen.

Wir definieren weiter:

Definition 8. Sei R ein H. O. f ist ein Erweiterungselement von R , und f^* wird f durch R zugeordnet, wenn für alle sinnvollen Rg $(f, Rg) = (f^*, g)$ gilt. f, f^* gehören zur $+, 0, -$ -Klasse, je nachdem ob $\Im(f^*, f) >, =, < 0$ ist.

Satz 11. f^* ist durch f eindeutig bestimmt (vgl. Def. 8); wenn Rf Sinn hat, so ist f Erweiterungselement von der 0-Klasse, und zwar $f^* = Rf$. R ist dann und nur dann max., wenn es keine weiteren Erweiterungselemente von der 0-Klasse gibt.

Beweis. Wären f^*, f^{**} beide f zugeordnet, so müßte $(f^*, g) = (f^{**}, g)$, $(f^* - f^{**}, g) = 0$ sein für alle sinnvollen Rg , also auch für die ganze abg. lin. M. dieser g , §. Hieraus folgern wir wieder $f^* - f^{**} = 0$. Die zweite Behauptung ist klar. Bei der dritten ist die Hinreichendheit der Bedingung klar, die Notwendigkeit folgt daraus, daß jedes Erweiterungselement der 0-Klasse zur eig. Forts. von R benutzt werden könnte.

Die durch Satz 11 gegebene Charakterisierung der Max. legt nun die folgende Begriffsbildung nahe.

Definition 9. Ein H. O. R heiße hypermaxima! (kurz: hypermax.), wenn es keine anderen Erweiterungselemente von R gibt als die, für welche R Sinn hat. (D. h. wenn R max. ist und die $+-$ und $-$ -Klassen leer sind.)

Die Wichtigkeit dieses Begriffes³¹⁾ wird sich noch zeigen, vgl. z. B. Satz 36.

³¹⁾ Er stammt von Herrn E. Schmidt, vgl. Anm. ²³⁾. Bei Beschränkung auf den reellen Hilbertschen Raum tritt die Unterscheidung max. — hypermax. nicht auf.

Wir beweisen noch einen Satz über die Stetigkeit von lin. O.

Satz 12. *Für einen lin. O. R ist eine jede der folgenden Aussagen mit der Stetigkeit gleichwertig.*

$$\alpha) \quad |Rf| \leq C \cdot |f|, \quad \beta) \quad |(f, Rg)| \leq C \cdot |f| \cdot |g|.$$

Wenn R ein H. O. ist, so kommt noch hinzu:

$$\gamma) \quad |(f, Rf)| \leq C \cdot |f|^2 \quad ^{32}).$$

Bemerkung. Nach Hilbert sagen wir bei lin. H. O. für stetig auch beschränkt. Bei nicht-lin. O. unterlassen wir es, die Stetigkeit zu untersuchen. (Die Hilbertsche Definition ist β .)

Beweis. Wegen der Lin. folgt aus α) $|Rf - Rg| = |R(f - g)| \leq C|f - g|$, d. h. die Stetigkeit. In β) brauchen wir nur $f = Rg$ zu setzen, um $|Rg|^2 \leq C \cdot |Rg| \cdot |g|$, $|Rg| \leq C \cdot |g|$, d. h. α), zu gewinnen. Wenn R stetig ist, so gibt es wegen $R0 = 0$ (Lin.!) ein $\varepsilon > 0$, so daß $|f| \leq \varepsilon$ $|Rf| \leq 1$ nach sich zieht. Daher ist stets (Lin.!) $|Rf| \leq \frac{1}{\varepsilon} \cdot |f|$, also $|(f, Rg)| \leq |f| \cdot |Rg| \leq \frac{1}{\varepsilon} \cdot |f| \cdot |g|$, d. h. β) gilt. Somit sind α), β) der Stetigkeit wirklich gleichwertig.

γ) besagt offenbar weniger als β), also genügt es β) für H. O. aus γ) zu folgern. Das geschieht ähnlich wie bei Satz 1. Wir setzen für f $\frac{f+g}{2}$ und $\frac{f-g}{2}$ ein und subtrahieren: dann wird

$$\Re(f, Rg) \leq C \cdot \left(\frac{1}{2} |f|^2 + \frac{1}{2} |g|^2 \right).$$

Indem wir (wie bei Satz 1) f, g durch $af, \frac{1}{a}g$ ($a > 0$) ersetzen und das Minimum der rechten Seite nehmen, sodann f, g durch $\vartheta f, g$ ($|\vartheta| = 1$), und das Maximum der linken bilden, erhalten wir $|(f, Rg)| \leq C \cdot |f| \cdot |g|$, d. h. β). Aber hier ist Rf, Rg als sinnvoll vorausgesetzt. Von der Sinnvollheit von Rf befreien wir uns: in β) kommt Rf nicht vor, es gilt für alle f , für die dies Sinn hat, also in einer überall dichten Menge (weil R ein lin. H. O. ist); da beide Seiten stetig sind, gilt es also für alle f überhaupt.

Für lin. H. O. haben wir als Bedingung der Beschränktheit:

$$-C \cdot |f|^2 \leq (f, Rf) \leq C \cdot |f|^2$$

(Satz 12 γ)). Dies zerlegen wir wie üblich:

³²⁾ Für einen H. O. muß jedes (f, Rf) reell sein, da ja $(f, Rf) = (Rf, f) = \overline{(f, Rf)}$ gilt.

Definition 10. Ein lin. H. O. R heie nach oben bzw. unten halb-beschrnkt, wenn stets

$$(f, Rf) \leq C \cdot |f|^2 \quad \text{bzw.} \quad \geq -C \cdot |f|^2$$

gilt. (Beides zusammen ist nach Satz 12 γ) Beschrntheit.) Wenn insbesondere $(f, Rf) \geq 0$ gilt, so heie er definit.

III. Lineare Mannigfaltigkeiten und Projektionsoperatoren.

Definition 11. $\mathfrak{M}, \mathfrak{N}$ seien zwei, nicht notwendig abg., lin. M. Dann sind $\mathfrak{M} + \mathfrak{N}$ und $\mathfrak{M} \times \mathfrak{N}$, die von ihrer Vereinigungsmenge aufgespannte lin. M. und ihr Durchschnitt, wieder lin. M. Wenn insbesondere $\mathfrak{M}, \mathfrak{N}$ abg. lin. M. sind, so ist es $\mathfrak{M} \times \mathfrak{N}$ offenbar ebenfalls. (Ob $\mathfrak{M} + \mathfrak{N}$ auch abg. ist, ist ungewi; fr zueinander orth. $\mathfrak{M}, \mathfrak{N}$ folgt es aus Satz 17.) Im folgenden sollen $\mathfrak{M}, \mathfrak{N}$ als abg. lin. M. vorausgesetzt werden. Wenn $\mathfrak{M}$ Teilmenge von $\mathfrak{N}$ ist, so ist $\mathfrak{N} - \mathfrak{M}$ ihre Differenz, d. h. die Menge aller Elemente von $\mathfrak{N}$, die zu allen Elementen von $\mathfrak{M}$ orth. sind, auch eine abg. lin. M. $\mathfrak{N} - \mathfrak{M}$ heie auch die Komplementre von $\mathfrak{M}$.

Satz 13. *Sei $\mathfrak{M}$ eine abg. lin. M. Jedes f kann auf eine und nur eine Art in $f = g + h$, g von $\mathfrak{M}$, h von $\mathfrak{N} - \mathfrak{M}$ zerlegt werden.*

Bemerkung. g heit dann die Projektion von f in $\mathfrak{M}$, die Zuordnung von g zu f ist offenbar ein berall sinnvoller Operator, wir nennen ihn den Projektionsoperator von $\mathfrak{M}$, $P_{\mathfrak{M}}$ ($\mathfrak{M}$ ist eine abg. lin. M.!).

Beweis. Die Eindeutigkeit ist klar: Aus $g_1 + h_1 = g_2 + h_2$ (g_1, g_2 von $\mathfrak{M}$, h_1, h_2 von $\mathfrak{N} - \mathfrak{M}$) folgt $g_1 - g_2 = h_2 - h_1$, dieses gehrt also zu $\mathfrak{M}$ wie zu $\mathfrak{N} - \mathfrak{M}$, ist somit zu sich selbst orth., also $= 0$; d. h. $g_1 = g_2$, $h_1 = h_2$. Es bleibt brig, die Existenz zu beweisen.

Da $\mathfrak{N}$ separabel ist, gibt es eine in $\mathfrak{M}$ berall dichte Folge $f_1, f_2, \dots$ ³³⁾, diese spannt also die abg. lin. M. $\mathfrak{M}$ auf. Nach Satz 8 gibt es also ein norm. orth. System $\varphi_1, \varphi_2, \dots$, welches dasselbe tut. Nach Satz 6 konvergiert $g = \sum_{n=1}^{\infty} (f, \varphi_n) \cdot \varphi_n$, und es gehrt mit den $\varphi_1, \varphi_2, \dots$ zu $\mathfrak{M}$, und $h = f - g$ ist zu allen $\varphi_1, \varphi_2, \dots$ orth. Also auch zu ihrer abg. lin. M., $\mathfrak{M} -$ d. h. h gehrt zu $\mathfrak{N} - \mathfrak{M}$. Damit ist alles bewiesen.

Satz 14. *Der Projektionsoperator $P_{\mathfrak{M}}$ ist ein berall sinnvoller max. H. O., es ist $P_{\mathfrak{M}}^2 = P_{\mathfrak{M}}$.*

Beweis. Da $P_{\mathfrak{M}}$ berall sinnvoll ist, wissen wir. Er ist ein H. O., wenn wir $(f, P_{\mathfrak{M}} g) = (P_{\mathfrak{M}} f, g)$ beweisen knnen, d. h. fr $f = h + k$,

³³⁾ Jede Teilmenge einer separablen topologischen Menge ist wieder separabel (siehe z. B. Hausdorff, Einfhrung in die Mengenlehre, Leipzig-Berlin 1914).

$g = h' + k'$ (h, h' von $\mathfrak{M}$, k, k' von $\mathfrak{S} - \mathfrak{M}$) (f, h') = (h, g). Nun ist (k, h') = 0, (h, k') = 0, also

$$(f, h') = (h + k, h') = (h, h') = (h, h' + k') = (h, g).$$

Wenn f zu $\mathfrak{M}$ gehört, ist die Zerlegung von Satz 13 $f = f + 0$, also $P_{\mathfrak{M}} f = f$. Da $P_{\mathfrak{M}} f$ zu $\mathfrak{M}$ gehört, ist also $P_{\mathfrak{M}} (P_{\mathfrak{M}} f) = P_{\mathfrak{M}} f$, d. h. $P_{\mathfrak{M}}^2 = P_{\mathfrak{M}}$. Da $P_{\mathfrak{M}}$ überall sinnvoll ist, muß es max. sein. —

Die max. H. O. E mit der Eigenschaft $E^2 = E$ haben aber für uns ein selbständiges Interesse (sie entsprechen den „Einzelformen“ bei Hilbert), wir untersuchen sie darum zunächst für sich.

Satz 15. Zu jedem max. H. O. E mit der Eigenschaft $E^2 = E$ ³⁴⁾ gibt es eine abg. lin. M. $\mathfrak{M}$, so daß $P_{\mathfrak{M}} = E$ ist. $\mathfrak{M}$ ist der Wertevorrat von E .

Beweis. Da $\mathfrak{M}$ offenbar der Wertevorrat von $P_{\mathfrak{M}}$ ist, genügt es zu zeigen, daß $P_{\mathfrak{M}} = E$ ist, wenn $\mathfrak{M}$ der Wertevorrat von E ist.

Wegen $E^2 = E$ ist, wie man leicht einsieht, $\mathfrak{M}$ die Menge aller f mit $Ef = f$, also eine abg. lin. M. (da E max. ist, ist es auch lin. abg.!). Wenn Ef Sinn hat, so ist $f - Ef$ zu jedem Element von $\mathfrak{M}$, d. h. jedem Eq , orth.:

$$(Eq, f - Ef) = (q, Ef - E^2 f) = 0;$$

also gehört es zu $\mathfrak{S} - \mathfrak{M}$. Daher ist $P_{\mathfrak{M}} f = Ef$, also $P_{\mathfrak{M}}$ Forts. von E . Da E max. ist, ist folglich $P_{\mathfrak{M}} = E$.

Definition 12. Einen Operator, der den Bedingungen von Satz 15 genügt, nennen wir Projektionsoperator (kurz: P. O., nach Satz 15 sind die P. O. mit den $P_{\mathfrak{M}}$ identisch und den abg. lin. M. $\mathfrak{M}$ ein-eindeutig zugeordnet). Eine abg. lin. M. $\mathfrak{M}$ und ihren P. O. $P_{\mathfrak{M}}$ nennen wir zueinandergehörig. Man sieht sofort: zu den abg. lin. M. (0) ³⁵⁾ und $\mathfrak{S}$ gehören die P. O. 0 bzw. 1 ³⁶⁾, ferner ist $P_{\mathfrak{S} - \mathfrak{M}} = 1 - P_{\mathfrak{M}}$. (Also sind 0, 1 P. O., und mit E ist es auch $1 - E$.)

Satz 16. Ein P. O. E ist stets beschränkt und definit, es ist nämlich

$$(f, Ef) = |Ef|^2, \quad |Ef| \leq |f|.$$

Beweis. Es ist $(f, Ef) = (f, E^2 f) = (Ef, Ef) = |Ef|^2$, also ist $(f, Ef) \geq 0$. Wenden wir dies auf den P. O. $1 - E$ an, so wird: $(f, (1 - E)f) \geq 0$, $(f, Ef) \leq (f, f) = |f|^2$, also $|Ef|^2 \leq |f|^2$, $|Ef| \leq |f|$. Damit sind die Formeln der zweiten Zeile bewiesen. Die erste Zeile folgt aus ihnen.

³⁴⁾ D. h. wenn Ef sinnvoll ist, so ist es auch $E(Ef)$, und zwar $= Ef$.

³⁵⁾ (0) ist die Menge, deren einziges Element die 0 (aus $\mathfrak{S}$) ist. Die leere Menge sehen wir nicht als lin. M. an.

³⁶⁾ 0, 1 sind zwei überall sinnvolle Operatoren, durch $0f = 0$, $1f = f$ definiert.

Satz 17. $\mathfrak{M}, \mathfrak{N}$ seien zwei abg. lin. M ., E, F die zugehörigen P. O. EF ist dann und nur dann ein P. O., wenn $EF = FE$ ist, d. h. E, F vertauschbar; dann nennen wir auch $\mathfrak{M}, \mathfrak{N}$ vertauschbar, EF ist der P. O. von $\mathfrak{M} \times \mathfrak{N}$. $E + F$ ist dann und nur dann ein P. O., wenn $EF = 0$ ist (oder auch: wenn $FE = 0$ ist); dies bedeutet, daß ganz $\mathfrak{M}$ auf ganz $\mathfrak{N}$ orth. steht, wir nennen dann auch E, F orth., $E + F$ ist der P. O. von $\mathfrak{M} + \mathfrak{N}$. $E - F$ ist dann und nur dann ein P. O., wenn $EF = F$ ist (oder auch: wenn $FE = F$ ist); dies bedeutet, daß $\mathfrak{N}$ Teilmenge von $\mathfrak{M}$ ist, wir nennen dann auch F Teil von E , $E - F$ ist der P. O. von $\mathfrak{M} - \mathfrak{N}$.³⁷⁾

Beweis. EF ist überall sinnvoll, also ein max. H. O., wenn es der Symmetrie-Bedingung $(EFf, g) = (f, EFg)$ genügt. Nun ist

$$(f, EFg) = (Ef, Fg) = (FEf, g), \quad (EFf, g) - (f, EFg) = (\{EF - FE\}f, g),$$

damit dies für alle g verschwinde, muß $(EF - FE)f = 0$ sein (für alle f), also $EF - FE = 0$. Unsere Bedingung ist also bereits für den max. H. O.-Charakter von EF notwendig und hinreichend. Und da sie

$$(EF)^2 = EFFE = EEFF = E^2 F^2 = EF$$

zur Folge hat, auch für den P. O.-Charakter.

Für $E + F$ ist der H. O.-Charakter klar, ebenso die Max. (es ist überall sinnvoll!), es genügt also $(E + F)^2 = E + F$ zu untersuchen. Wegen

$$(E + F)^2 = E^2 + F^2 + EF + FE = (E + F) + (EF + FE)$$

bedeutet es $EF + FE = 0$. Hieraus folgt aber $EF = 0$, denn man multipliziere links mit E : $EF + EFE = 0$, und dieses rechts mit E : $2 \cdot EFE = 0$ beides zusammen ergibt $EF = 0$ ³⁸⁾. Ist umgekehrt $EF = 0$, so ist es ein P. O., also $FE = EF = 0$, $EF + FE = 0$. Somit ist $EF = 0$ wirklich charakteristisch, Vertauschen von E, F zeigt, daß es $FE = 0$ gleichfalls ist.

$E - F$ ist dann und nur dann ein P. O., wenn $1 - (E - F) = (1 - E) + F$ einer ist, da $1 - E$ und F solche sind, bedeutet dies: $(1 - E)F = 0$, $F - EF = 0$, oder auch: $F(1 - E) = 0$, $F - FE = 0$.

Es bleibt noch übrig, die Aussagen über $\mathfrak{M}, \mathfrak{N}, \mathfrak{M} \times \mathfrak{N}, \mathfrak{M} \pm \mathfrak{N}$ zu beweisen. Sei erstens $EF = FE$. Jedes $EFf = FEFf$ gehört zu $\mathfrak{M}$ und $\mathfrak{N}$, also zu $\mathfrak{M} \times \mathfrak{N}$; wenn aber f zu $\mathfrak{M} \times \mathfrak{N}$, d. h. $\mathfrak{M}$ und $\mathfrak{N}$, gehört, so ist $EFf = Ef = f$ — d. h. $\mathfrak{M} \times \mathfrak{N}$ ist der Wertevorrat von EF , also dieses sein P. O. Sei zweitens $EF = 0$ (also $FE = 0$). Wenn f zu $\mathfrak{M}$, g zu $\mathfrak{N}$ gehört, so ist $(f, g) = (Ef, Fg) = (f, EFg) = 0$, d. h. $\mathfrak{M}$ zu $\mathfrak{N}$ orth., ebenso klar ist die Umkehrung. Für jedes f gehört Ef zu $\mathfrak{M}$, Ff zu $\mathfrak{N}$,

³⁷⁾ Da es sich um überall sinnvolle Operatoren handelt, dürfen wir ruhig $EF, E \pm F$ bilden; sonst wären die Definitionsbereiche dieser Operatoren noch genauer zu erörtern.

³⁸⁾ Dieser Kunstgriff stammt von Herrn E. Schmidt.

also $(E + F)f$ zu $\mathfrak{M} + \mathfrak{N}$; wenn f zu $\mathfrak{M} + \mathfrak{N}$ gehört, so ist es $= g + h$, g von $\mathfrak{M}$, h von $\mathfrak{N}$, also

$(E + F)f = (E + F)(g + h) = Eg + Fg + Eh + Fh = g + 0 + 0 + h = f$
 — d. h. $\mathfrak{M} + \mathfrak{N}$ ist der Wertevorrat von $E + F$, also eine abg. lin. M., und dieses sein P. O. (Satz 15). Sei drittens $EF = F$ (also $FE = F$). Dies bedeutet $EFf = Ff$ für alle f , d. h. $Eg = g$ für alle g von $\mathfrak{N}$, d. h. (da dies die Definitionsgleichung von $\mathfrak{M}$ ist, vgl. den Beweis von Satz 15) daß $\mathfrak{N}$ Teilmenge von $\mathfrak{M}$ ist. Wegen $E - F = E - EF = E(1 - F)$ ist $E - F$ der P. O. von $\mathfrak{M} \times (\mathfrak{S} - \mathfrak{N}) = \mathfrak{M} - \mathfrak{N}$.

Damit sind alle unsere Behauptungen bewiesen.

Satz 18. E ist dann und nur dann Teil von F , wenn stets

$$|Ef| \leq |Ff|$$

gilt. $E_1 + \dots + E_p$ ist dann und nur dann P. O. ($E_1, \dots, E_p$ P. O.), wenn

$$E_m E_n = 0 \quad (m \neq n; m, n = 1, \dots, p)$$

gilt.

Beweis. Wenn E Teil von F ist, so ist $E = EF$, also $|Ef| = |EFF| \leq |Ff|$. Gilt dies umgekehrt für alle f , so hat $f = Ef$ die Konsequenz

$$|Ff| \geq |Ef| = |f|, \quad (f, Ff) = |Ff|^2 \geq |f|^2 = (f, f),$$

$$|(1 - F)f|^2 = (f, (1 - F)f) \leq 0, \quad (1 - F)f = 0,$$

also $f = Ff$. D. h. wenn E, F die P. O. von $\mathfrak{M}, \mathfrak{N}$ sind, so ist $\mathfrak{M}$ Teil von $\mathfrak{N}$ — also E Teil von F .

Daß $E_m E_n = 0$ ($m \neq n$) für den P. O.-Charakter von $E_1 + \dots + E_p$ hinreicht, folgt durch wiederholte Anwendung von Satz 17; seine Notwendigkeit erkennen wir so. Es ist ($m \neq n$)

$$|E_m f|^2 + |E_n f|^2 \leq |E_1 f|^2 + \dots + |E_p f|^2 = (f, E_1 f) + \dots + (f, E_p f)$$

$$= (f, (E_1 + \dots + E_p)f) \leq |f|^2,$$

also für $E_n f = f$ $|E_m f|^2 \leq 0$, $E_m f = 0$. Da für $f = E_n g$ das erstere zutrifft, ist stets $E_m E_n g = 0$, d. h. $E_m E_n = 0$, wie behauptet wurde. —

Zum Schluß noch ein Satz über die Konvergenz von P. O.

Satz 19. $E_1, E_2, \dots$ sei eine Folge von P. O., und entweder E_m Teil von E_n für alle $m < n$, oder E_n Teil von E_m für alle $m < n$. Dann gibt es einen P. O. E , so daß für jedes f $E_n f \rightarrow Ef$ gilt.

Beweis. Wenn wir $E_1, E_2, \dots, E$ durch $1 - E_1, 1 - E_2, \dots, 1 - E$ ersetzen, geht der erste Unterfall in den zweiten über³⁹⁾, es genügt daher etwa den ersteren zu betrachten.

³⁹⁾ Denn wenn E Teil von F ist, so ist $1 - F$ Teil von $1 - E$, z. B. darum, weil $(1 - E)(1 - F) = 1 - E - F + EF = 1 - F$ ist.

Mit wachsendem n wächst $|E_n f|^2$ (Satz 18), und es bleibt $\leq |f|^2$ (Satz 16), daher existiert zu jedem $\varepsilon > 0$ ein $N = N(\varepsilon)$, so daß für jedes $m, n \geq N$ $|E_n f|^2 - |E_m f|^2 \leq \varepsilon$ bleibt. Sei noch $m < n$, dann ist $E_n - E_m$ ein P. O., also

$$|E_n f - E_m f|^2 = (f, (E_n - E_m) f) = (f, E_n f) - (f, E_m f) = |E_n f|^2 - |E_m f|^2 \leq \varepsilon.$$

Somit hat die Folge $E_1 f, E_2 f, \dots$ einen Limes, er heiße $E f$.

E ist ein überall sinnvoller Operator. Da

$$(f, E g) = \lim (f, E_n g) = \lim (E_n f, g) = (E f, g)$$

ist, ist es ein max. H. O. Wir brauchen daher nur noch $E^2 = E$ zu beweisen. Es ist

$$\begin{aligned} (f, E^2 g) &= (E f, E g) = \lim (E_n f, E_n g) = \lim (f, E_n^2 g) = \lim (f, E_n g) \\ &= (f, E g), \end{aligned}$$

also in der Tat $E^2 = E$.

IV. Reduzibilität von Operatoren.

Bei endlichvieldeimensionalen Matrizen versteht man unter Reduzibilität bekanntlich die Möglichkeit, die Matrix unitär auf eine solche Form zu transformieren, daß in ihr (die etwa $N = M_1 + M_2$ -dimensional sein mag) nur die gemeinsamen Felder der M_1 ersten Zeilen und Spalten und die gemeinsamen Felder der M_2 letzten Zeilen und Spalten anders als mit Nullen besetzt sein können. Wenn die Matrix als lineare Transformation aufgefaßt wird, so heißt dies: der Raum ist in zwei (zueinander orth.) lin. M. von M_1 - bzw. M_2 -Dimensionen zerlegbar, deren jede durch die Transformation in ein Teil von sich übergeführt wird.

Bei den Operatoren in § analogisieren wir dies folgendermaßen, wobei wir uns auf abg. lin. Operatoren beschränken:

Definition 13. Ein abg. lin. Operator R heiße durch die abg. lin. M. $\mathfrak{M}$ reduziert, wenn aus der Sinnvollheit von $R f$ auch die von $R P_{\mathfrak{M}} f$ folgt, und zwar dieses $= P_{\mathfrak{M}} R f$ ist⁴⁰⁾.

Durch (0) und § wird offenbar jedes R reduziert, und wenn $\mathfrak{M}$ das R reduziert, so tut es auch $\mathfrak{S} - \mathfrak{M}$ (wegen $P_{\mathfrak{S}-\mathfrak{M}} = 1 - P_{\mathfrak{M}}$). Irreduzibel heiße R , wenn es nur durch (0) und § reduziert wird.

Wenn R durch $\mathfrak{M}$ reduziert wird, so gehört für jedes f von $\mathfrak{M}$ offenbar auch $R f$ zu $\mathfrak{M}$ ($R f = R P_{\mathfrak{M}} f = P_{\mathfrak{M}} R f$), also erzeugt R auch einen Operator in $\mathfrak{M}$ (d. h. eine Funktion, deren Definitionsbereich und Wertevorrat Teilmengen von $\mathfrak{M}$ sind). Diesen Operator nennen wir den in $\mathfrak{M}$ liegenden Teil von R .

⁴⁰⁾ Dies kann also auch als Vertauschbarkeit von R mit $P_{\mathfrak{M}}$ aufgefaßt werden.

Wir beweisen nun zwei Sätze über das Reduzieren von Operatoren.

Satz 20. *R sei ein abg. lin. Operator, $\mathfrak{M}$ eine abg. lin. M. Damit $\mathfrak{M}$ R reduziere, ist notwendig, daß mit Rf auch $RP_{\mathfrak{M}}f$ stets sinnvoll sei, und mit f auch Rf (falls es sinnvoll ist) stets zu $\mathfrak{M}$ gehöre. Wenn R ein H.O. ist, so ist dies auch hinreichend⁴¹⁾.*

Beweis. Die Notwendigkeit haben wir schon vorher erkannt, es bleibt zu zeigen, daß dies bei H.O. R auch hinreichend ist, d. h. $RP_{\mathfrak{M}}f = P_{\mathfrak{M}}Rf$ zu beweisen. Da $RP_{\mathfrak{M}}f$ nach Annahme zu $\mathfrak{M}$ gehört, ist die Zugehörigkeit von $Rf - RP_{\mathfrak{M}}f = R(f - P_{\mathfrak{M}}f)$ zu $\mathfrak{S} - \mathfrak{M}$ zu zeigen, oder da $f - P_{\mathfrak{M}}f$ zu $\mathfrak{S} - \mathfrak{M}$ gehört: wenn g zu $\mathfrak{S} - \mathfrak{M}$ gehört, so gehört auch Rg dazu (wenn es Sinn hat).

Sei Rf sinnvoll, dann gehört $RP_{\mathfrak{M}}f$ zu $\mathfrak{M}$, also ist

$$(f, P_{\mathfrak{M}}Rg) = (P_{\mathfrak{M}}f, Rg) = (RP_{\mathfrak{M}}f, g) = 0.$$

Da diese f überall dicht liegen, ist $P_{\mathfrak{M}}Rg$ zu allem orth., also $= 0$ — daher gehört Rg zu $\mathfrak{S} - \mathfrak{M}$.

Satz 21. $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ sei eine (eventuell abbrechende) Folge von paarweise orth. at. lin. M., die zusammen die abg. lin. M. $\mathfrak{S}$ aufspannen. R sei ein abg. lin. Operator, der durch jede von ihnen reduziert wird, eine in $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ gelegenen Teile seien bzw. $R_1, R_2, \dots$.

Dann wird R durch $R_1, R_2, \dots$ bestimmt, und zwar folgendermaßen: Rf hat dann und nur dann Sinn, wenn $f = f_1 + f_2 + \dots$ ist ($f_1, f_2, \dots$ aus bzw. $\mathfrak{M}_1, \mathfrak{M}_2, \dots$), alle $R_1f_1, R_2f_2, \dots$ Sinn haben, beide Reihen $f_1 + f_2 + \dots, R_1f_1 + R_2f_2 + \dots$ konvergieren (wenn die Folge $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ abbricht, so fällt diese Konvergenzannahme natürlich fort), und zwar ist dann $Rf = R_1f_1 + R_2f_2 + \dots$.

Beweis. Daß R an den angeführten Stellen Sinn hat, und zwar die angegebenen Werte annimmt, das folgt sofort aus der Definition von $R_1, R_2, \dots$ sowie dem abg. lin. und durch $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ reduzierten Charakter von R . Wir müssen noch beweisen: wenn Rf Sinn hat, so hat f die obige Form.

Wir setzen $f_p = P_{\mathfrak{M}_p}f$ ($p = 1, 2, \dots$), f_p liegt dann in $\mathfrak{M}_p$, und die $P_{\mathfrak{M}_p}$ (wie die $\mathfrak{M}_p$) sind paarweise orth. Auch ist $R_p f_p = R f_p = R P_{\mathfrak{M}_p} f = P_{\mathfrak{M}_p} R f$. Aus der Orth. der $P_{\mathfrak{M}_p}$ folgt, daß die $P_{\mathfrak{M}_1}, P_{\mathfrak{M}_1} + P_{\mathfrak{M}_2}, P_{\mathfrak{M}_1} + P_{\mathfrak{M}_2} + P_{\mathfrak{M}_3}, \dots$ eine (natürlich wachsende) Folge von P.O. bilden, die nach Satz 19 einen P.O. E zum Limes haben. Aus Stetigkeitsgründen sind alle Teile von E , also die $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ Teile der zu E gehörigen abg. lin. M. $\mathfrak{M}$. Nach der Annahme über die $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ ist daher $\mathfrak{M} = \mathfrak{S}$, $E = 1$.

⁴¹⁾ In endlichvioldimensionalen Räumen ist dies auch bei unitären Operatoren hinreichend, in $\mathfrak{S}$ nicht; das ist aber hier belanglos.

Also sind die Reihen $f_1 + f_2 + \dots$, $R_1 f_1 + R_2 f_2 + \dots$ konvergent, und sie haben die bzw. Summen $1f = f$, $1Rf = Rf$. Damit ist alles bewiesen.

V. Die Cayleysche Transformation.

Nach diesen Vorbereitungen können wir unsere eigentliche Aufgabe in Angriff nehmen, die genauere Untersuchung der verschiedenen Arten von H. O. Die Grundlage dieser Überlegungen ist ein Kunstgriff, der eine Übertragung der Cayleyschen Transformation der Algebra auf unser Gebiet ist.

Sei R ein abg. lin. H. O. und $\mathfrak{A}$ sein Definitionsbereich. $\mathfrak{A}$ ist eine lin. M. und überall dicht (also nur dann abg., wenn es $= \mathfrak{H}$, d. h. R überall sinnvoll ist, was keineswegs vorausgesetzt wird). Wir betrachten die beiden Operatoren $R \pm i1$, die auch in $\mathfrak{A}$ definiert sind. Man berechnet mühelos:

$$\begin{aligned} ((R \pm i1)f, (R \pm i1)g) &= (Rf, Rg) \pm (Rf, ig) \pm (if, Rg) + (f, g) \\ &= (Rf, Rg) + (f, g), \end{aligned}$$

also insbesondere für $f = g$

$$|(R \pm i1)f| = |(R - i1)f| = \sqrt{|Rf|^2 + |f|^2}.$$

Somit hat $f \neq 0$ $(R + i1)f \neq 0$, $(R - i1)f \neq 0$ zur Folge, und wegen der Lin. $f \neq g$ die Folge $(R + i1)f \neq (R + i1)g$, $(R - i1)f \neq (R - i1)g$. Wenn $\mathfrak{A}$ durch $R \pm i1$ auf $\mathfrak{E}$ bzw. $\mathfrak{F}$ abgebildet wird, so sind die Abbildungen demnach alle ein-eindeutig. Insbesondere sei die Abbildung von $\mathfrak{E}$ auf $\mathfrak{F}$ U . Da $R \pm i1$, wie R , abg. lin. sind, gilt dasselbe von U ; weiter folgt aus der Gleichung von vorhin $|Uf| = |f|$. Nach Def. 6 ist daher U ein längentreuer Operator, da $\mathfrak{E}$, $\mathfrak{F}$ sein Definitionsbereich bzw. Wertevorrat ist, sind nach Satz 10 beide abg. lin. M.

Satz 22. *Sei R ein abg. lin. H. O. und $\mathfrak{A}$ sein Definitionsbereich; die Operatoren $R \pm i1$ mögen $\mathfrak{A}$ auf $\mathfrak{E}$ bzw. $\mathfrak{F}$ abbilden. Diese Abbildungen von $\mathfrak{A}$, $\mathfrak{E}$, $\mathfrak{F}$ aufeinander sind ein-eindeutig, und insbesondere ist die Abbildung von $\mathfrak{E}$ auf $\mathfrak{F}$ ein längentreuer Operator U . $\mathfrak{E}$, $\mathfrak{F}$ sind beide abg. lin. M.*

Beweis. Wurde soeben erbracht.

Definition 14. Die Operatoren R , U aus Satz 22 nennen wir Cayleysche Transformierte voneinander, R ist immer ein abg. lin. H. O., U immer längentreu⁴²⁾.

⁴²⁾ Zunächst wissen wir nur, daß jedes abg. lin. H. O. R genau eine Cayleysche Transformierte U (längentreu) hat. Der Gültigkeitsbereich der Umkehrung wird aus Satz 24 hervorgehen.

Satz 23. *Ein längentreuer Operator U ist dann und nur dann Cayleysche Transformierte des abg. lin. H. O. R , wenn*

α) $\mathfrak{A}$ die Menge aller $\varphi - U\varphi$, φ von $\mathfrak{E}$, ist,

β) für alle φ von $\mathfrak{E}$ $R(\varphi - U\varphi) = i(\varphi + U\varphi)$ ist.

Beweis. Sei U die Cayleysche Transformierte von R . Wenn f zu $\mathfrak{A}$ gehört, so gehört $(R + i1)f = \varphi$ zu $\mathfrak{E}$, es ist $U\varphi = (R - i1)f$, also $\varphi - U\varphi = 2if$, $f = \frac{\varphi}{2i} - U\frac{\varphi}{2i}$. Gehört umgekehrt φ zu $\mathfrak{E}$, so ist $(R + i1)f = \varphi$ für ein f von $\mathfrak{A}$, also $(R - i1)f = U\varphi$, $\varphi - U\varphi = 2if$, d. h. $\varphi - U\varphi$ gehört zu $\mathfrak{A}$. Und schließlich ist $\varphi + U\varphi = 2Rf$, also $R(\varphi - U\varphi) = 2iRf = i(\varphi + U\varphi)$. Also sind α), β) erfüllt.

Genüge umgekehrt $U \alpha$), β). Wenn f zu $\mathfrak{A}$ gehört, so ist $f = \varphi - U\varphi$, φ von $\mathfrak{E}$ (α)), $Rf = i(\varphi + U\varphi)$ (β)), also $(R + i1)f = 2i\varphi$, $(R - i1)f = 2iU\varphi$. Somit gehört $(R + i1)f$ zu $\mathfrak{E}$, und es ist $U(R + i1)f = (R - i1)f$. Wenn wir dagegen von einem φ von $\mathfrak{E}$ ausgehen, so gehört $f = \varphi - U\varphi$ zu $\mathfrak{A}$, also ist $Rf = i(\varphi + U\varphi)$, $2i\varphi = (R + i1)f$, $\varphi = (R + i1)\frac{f}{2i}$. Somit ist $\mathfrak{E}$ die Menge aller $(R + i1)f$, f von $\mathfrak{A}$, und $U(R + i1)f = (R - i1)f$, d. h. U wirklich die Cayleysche Transformierte von R .

Satz 24. *Zum längentreuen Operator U gibt es dann und nur dann einen abg. lin. H. O. R , von dem er Cayleysche Transformierte ist, wenn die Menge aller $\varphi - U\varphi$, φ von $\mathfrak{E}$ ($\mathfrak{E}$ ist der Definitionsbereich von U), überall dicht ist. Und zwar ist dann R durch U eindeutig bestimmt.*

Beweis. Die Notwendigkeit ist klar nach α), Satz 23, da $\mathfrak{A}$ überall dicht ist. Die Hinreichendheit sowie die eindeutige Bestimmtheit von R durch U (nach α), β) in Satz 23) steht gleichfalls fest, sobald erkannt ist, daß β) für R keine Unmöglichkeit involviert, und daß das so definierte R ein abg. lin. H. O. ist.

Das erstere folgt daraus, daß für $\varphi \neq \psi$ $\varphi - U\varphi \neq \psi - U\psi$ ist, d. h. für $\varphi \neq 0$ $\varphi - U\varphi \neq 0$. In der Tat hat $\varphi = U\varphi$ die Folge, daß für alle χ von $\mathfrak{E}$

$$(\varphi, \chi) = (U\varphi, U\chi) = (\varphi, U\chi), \quad (\varphi, \chi - U\chi) = 0$$

ist, da aber die $\chi - U\chi$ überall dicht liegen, ist φ zu einer überall dichten Menge orth., also zu ganz $\mathfrak{H}$, also $\varphi = 0$. Das letztere beweisen wir so. Die abg. Lin. von R folgt aus der von U , die überall Dichtigkeit des Definitionsbereiches von R (d. h. der $\varphi - U\varphi$) wurde vorausgesetzt — es bleibt noch übrig $(f, Rg) = (Rf, g)$ zu beweisen, d. h. daß (f, Rg) beim Vertauschen von f und g in seine komplex-Konjugierte übergeht. Sei also $f = \varphi - U\varphi$, $g = \psi - U\psi$, dann ist

$$\begin{aligned}
 (f, Rg) &= (\varphi - U\varphi, i(\psi + U\psi)) \\
 &= -i\{(\varphi, \psi) - (U\varphi, \psi) + (\varphi, U\psi) - (U\varphi, U\psi)\} \\
 &= i((U\varphi, \psi) - \overline{(U\psi, \varphi)}) = i(U\varphi, \psi) + i\overline{(U\psi, \varphi)},
 \end{aligned}$$

was die Behauptung in Evidenz setzt.

Satz 25. *R, S seien abg. lin. H. O., U, V ihre Cayleyschen Transformierten. R ist dann und nur dann Forts. oder eig. Forts. von S , wenn U Forts. bzw. eig. Forts. von V ist. — $\mathfrak{M}$ sei eine abg. lin. M., es reduziert R dann und nur dann, wenn es U reduziert.*

Beweis. Man schließt jeweils mühelos mit Hilfe von Satz 22 von R auf U , und mit Hilfe von Satz 23 von U auf R .

In den vorhergehenden Sätzen haben wir die Beurteilung der Fortsetzbarkeitsverhältnisse von R auf die viel einfacheren von U zurückgeführt. Die Cayleyschen Transformierten aller R sind alle U des Satzes 5, und dabei ist zu beachten: wenn R Forts. von S sein soll, also U Forts. eines Satz 5 genügenden V , so braucht es nur längentreu zu sein⁴⁸). An Stelle der zunächst recht unübersichtlichen Forts. von R brauchen wir also bloß die einfach angebbaren längentreuen Forts. von V (das selbst längentreu ist) zu betrachten. Um uns über derartige Fragen zu orientieren, beweisen wir zwei einfache (anschaulich-geometrische) Sätze über längentreue Abbildungen und abg. lin. M.

Zunächst bemerken wir: unter der Dimensionszahl einer abg. lin. M. $\mathfrak{M}$ verstehen wir die obere Grenze aller (endlichen!) n , für die es in $\mathfrak{M}$ n lin. unabh. Elemente gibt — also eine der Zahlen $0, 1, 2, \dots, \infty$. (Eine wesentlich allgemeinere Definition der Dimensionszahl von Mengen in $\S$ werden wir im nächsten Kapitel, Def. 16, geben.)

Satz 26. *$\mathfrak{M}$ sei eine abg. lin. M., $\varphi_1, \dots, \varphi_n$ ein norm. orth. System, welches diese abg. lin. M. aufspannt (es gibt gewiß solche, vgl. den Beweis von Satz 13; es ist $n = 0, 1, 2, \dots, \infty$, im letzteren Falle bricht die Folge $\varphi_1, \varphi_2, \dots$ nicht ab), dann ist $\mathfrak{M}$ n -dimensional.*

Zwei abg. lin. M. $\mathfrak{M}, \mathfrak{N}$ können dann und nur dann Definitionsbereich bzw. Wertevorrat eines längentreuen Operators U sein, wenn sie gleichvieldeimensional sind. Sei $\varphi_1, \dots, \varphi_n$ ein $\mathfrak{M}$ aufspannendes norm. orth. System, dann entsprechen die $\mathfrak{N}$ aufspannenden norm. orth. Systeme $\psi_1, \dots, \psi_n$ diesen U ein-eindeutig: jedem wird durch die Bedingungen $U\varphi_1 = \psi_1, \dots, U\varphi_n = \psi_n$ genau ein solches U zugeordnet.

Beweis. $\varphi_1, \dots, \varphi_n$ seien ein norm. orth. System, das die abg. lin. M. $\mathfrak{M}$ aufspannt. Da $\varphi_1, \dots, \varphi_n$ lin. unabh. sind, hat $\mathfrak{M} \geq n$ Dimensionen,

⁴⁸) Denn die $\varphi - U\varphi$ umfassen die $\varphi - V\varphi$, liegen also wie diese überall dicht.

für $n = \infty$ sind wir also fertig, für endliches n ist zu zeigen: $n + 1$ Elemente von $\mathfrak{M}$ sind nie lin. unabh. In der Tat müssen sie die Form

$$f_\mu = a_{\mu 1} \varphi_1 + \dots + a_{\mu n} \varphi_n \quad (\mu = 1, \dots, n + 1)$$

haben, und zwischen den $n + 1$ n -dimensionalen Vektoren

$$a_{\mu 1}, \dots, a_{\mu n} \quad (\mu = 1, \dots, n + 1)$$

muß eine lineare Relation bestehen, also auch zwischen den f_μ .

Gehen wir zur zweiten und dritten Behauptung über. Sei $\mathfrak{M}$ der Definitionsbereich und $\mathfrak{N}$ der Wertevorrat des längentreuen Operators U . Wenn $\varphi_1, \dots, \varphi_n$ ein $\mathfrak{M}$ aufspannendes norm. orth. System ist, so ist (wegen der Längentreue von U) $\psi_1 = U\varphi_1, \dots, \psi_n = U\varphi_n$ ein $\mathfrak{N}$ aufspannendes norm. orth. System. Und $\mathfrak{M}$ wie $\mathfrak{N}$ haben n Dimensionen.

Umgekehrt seien $\mathfrak{M}, \mathfrak{N}$ zwei gleichviel- (etwa n -) dimensionale abg. lin. M. Sei $\varphi_1, \dots, \varphi_p$ bzw. $\psi_1, \dots, \psi_q$ je ein $\mathfrak{M}$ bzw. $\mathfrak{N}$ aufspannendes norm. orth. System, es muß $p = q = n$ sein. Die Reihen $\sum_{r=1}^n x_r \varphi_r$ und $\sum_{r=1}^n x_r \psi_r$ konvergieren unter denselben Umständen: bei endlichem n immer, bei $n = \infty$, falls $\sum_{r=1}^{\infty} |x_r|^2$ endlich ist (Satz 5) — also allenfalls, wenn $\sum_{r=1}^{\infty} |x_r|^2$ endlich ist; sie erschöpfen also $\mathfrak{M}$ bzw. $\mathfrak{N}$ ($\sum_{r=1}^n x_r \varphi_r$ hat mit φ_s das innere Produkt x_s , es bestimmt also $x_1, \dots, x_n$ eindeutig). Wir können also einen Operator U durch $U\left(\sum_{r=1}^n x_r \varphi_r\right) = \sum_{r=1}^n x_r \psi_r$ definieren, er hat den Definitionsbereich $\mathfrak{M}$ und den Wertevorrat $\mathfrak{N}$, und ist offenbar linear. Es ist $\left(f = \sum_{r=1}^n x_r \varphi_r\right)$

$$|f|^2 = |Uf|^2 = \sum_{r=1}^n |x_r|^2, \quad |f| = |Uf|,$$

also der Operator stetig, und weil sein Definitionsbereich abg. ist, ist er auch selbst abg. Also ist U längentreu.

Dabei hat Längentreue, also insbesondere abg. Lin., und $U\varphi_\mu = \psi_\mu$ ($\mu = 1, \dots, n$) $U\left(\sum_{r=1}^n x_r \varphi_r\right) = \sum_{r=1}^n x_r \psi_r$ zur Folge; daher ist unser U durch diese Eigenschaften eindeutig festgelegt. Damit sind alle Behauptungen unseres Satzes bewiesen.

Satz 27. *U sei ein längentreuer Operator, $\mathfrak{E}, \mathfrak{F}$ sein Definitionsbereich bzw. Wertevorrat (also zwei abg. lin. M.). Wir gewinnen alle längentreuen Forts. von U auf die folgende Weise:*

Man wähle zwei gleichvioldimensionale abg. lin. M. $\mathfrak{M}, \mathfrak{N}$, Teile von $\mathfrak{E} - \mathfrak{E}, \mathfrak{F} - \mathfrak{F}$, und einen längentreuen Operator U' mit diesem Definitionsbereich bzw. Wertevorrat. Diese bestimmen genau eine längentreue

*Forts. V von U , die so charakterisiert ist: Definitionsbereich bzw. Wertevorrat sind $\mathfrak{E} + \mathfrak{M}$ bzw. $\mathfrak{F} + \mathfrak{N}$, in $\mathfrak{E}$ gilt $Vf = Uf$, in $\mathfrak{M}$ gilt $Vf = U'f$. — Indem man $\mathfrak{M}, \mathfrak{N}, U'$ auf alle möglichen Weisen wählt, gewinnt man alle längentreuen *Forts. V von U , und zwar jede genau einmal.**

Beweis. Daß jede längentreue *Forts. V von U* so entsteht, ist klar: seien $\mathfrak{R}, \mathfrak{L}$ Definitionsbereich bzw. Wertevorrat von V , so genügt es, $\mathfrak{M} = \mathfrak{R} - \mathfrak{E}$, $\mathfrak{N} = \mathfrak{L} - \mathfrak{F}$, $U' =$ das auf $\mathfrak{M}$ eingeschränkte V zu setzen, und alles ist erfüllt. Auch daß unsere Bedingungen genau ein V (wenn überhaupt etwas) festlegen, ist klar: genügen ihr V' und V'' , so ist $V'f = V''f$ in ganz $\mathfrak{E}$ und ganz $\mathfrak{M}$, also wegen der Lin. auch im ganzen Definitionsbereiche $\mathfrak{E} + \mathfrak{M}$ — d. h. $V' = V''$. Es bleibt daher nur übrig zu zeigen, daß wirklich jedes System $\mathfrak{M}, \mathfrak{N}, U'$ mindestens ein längentreues V bestimmt.

Da $\mathfrak{M}$ bzw. $\mathfrak{N}$ Teil von $\mathfrak{H} - \mathfrak{E}$ bzw. $\mathfrak{H} - \mathfrak{F}$, also orth. zu $\mathfrak{E}$ bzw. $\mathfrak{F}$ ist, sind $\mathfrak{M}, \mathfrak{E}$ lin. unabh., also in $\mathfrak{E} + \mathfrak{M}$ die Zerlegung $f = g + h$, g von $\mathfrak{E}$, h von $\mathfrak{M}$, eindeutig. Wir definieren also $V(g + h) = Ug + U'h$ (g von $\mathfrak{E}$, h von $\mathfrak{M}$). Da g, h orth. sind, und ebenso $Ug, U'h$, und U, U' längentreu, so ist $|V(g + h)|^2 = |Ug + U'h|^2 = |Ug|^2 + |U'h|^2 = |g|^2 + |h|^2$, $|g + h|^2 = |g|^2 + |h|^2$. V ist in der abg. lin. M. $\mathfrak{E} + \mathfrak{M}$ definiert, offenbar lin., nach dem vorigen stetig — also abg., und nach dem vorigen längentreu. Außerdem, wie gewünscht, *Forts. von U und U' .*

Durch die Sätze 26, 27 übersehen wir alle längentreuen *Forts. eines längentreuen U* , also auch alle abg. lin. H. O. *Forts. eines abg. lin. H. O. R* ; wir können sie alle effektiv konstruieren. So sehen wir z. B.: R ist max., d. h. es hat keine eig. *Forts.*, wenn U keine hat — also wenn es unmöglich ist in Satz 26 $\mathfrak{M}, \mathfrak{N}$ von (0) verschieden (vgl. Anm. ³⁵) zu wählen. Dies ist offenbar nur für $\mathfrak{H} - \mathfrak{E} = (0)$ oder $\mathfrak{H} - \mathfrak{F} = (0)$ der Fall, d. h. für $\mathfrak{E} = \mathfrak{H}$ oder $\mathfrak{F} = \mathfrak{H}$. Die genauere Untersuchung dieser Fragen soll in den drei nächsten Kapiteln durchgeführt werden.

VI. Erweiterungselemente.

Wir wollen die bereits im letzten Kapitel zur Geltung gekommenen abg. lin. M. $\mathfrak{E}, \mathfrak{F}$ in ihrem Verhältnis zu den Begriffsbildungen von Def. 8 näher charakterisieren.

Für die Zwecke dieses Kapitels ist es bequem, die Terminologie des Restklassen-Kalküls der Algebra zu verwenden: wenn $\mathfrak{M}$ eine (nicht notwendig abg.!) lin. M. ist, so bedeute $f = g \dots \text{mod } \mathfrak{M}$, daß $f - g$ zu $\mathfrak{M}$ gehört. $R, U, \mathfrak{A}, \mathfrak{E}, \mathfrak{F}$ mögen die bisherige Bedeutung haben, wie z. B. in Satz 22.

Satz 28. $\mathfrak{H} - \mathfrak{E}$ bzw. $\mathfrak{H} - \mathfrak{F}$ ist die Menge aller erweiterbaren f , denen R i f bzw. $-i f$ zuordnet.

Beweis. Daß f^* dem f zugeordnet ist, bedeutet nach Satz 23

$$(f^*, \varphi - U\varphi) = (f, i(\varphi + U\varphi)) = (-if, \varphi + U\varphi).$$

$f^* = if$ bzw. $= -if$ bedeutet also, daß f (und mit ihm $f^* = \pm if$) auf alle φ (von $\mathfrak{E}$) bzw. alle $U\varphi$ (d. h. ganz $\mathfrak{F}$) orth. steht; d. h. daß es zu $\mathfrak{H} - \mathfrak{E}$ bzw. $\mathfrak{H} - \mathfrak{F}$ gehört. —

Daher gehört ganz $\mathfrak{H} - \mathfrak{E}$ bzw. ganz $\mathfrak{H} - \mathfrak{F}$ (mit Ausnahme der 0) zur $+-$ bzw. $-$ -Klasse der Erweiterungselemente.

Satz 29. $\mathfrak{A} + (\mathfrak{H} - \mathfrak{E}) + (\mathfrak{H} - \mathfrak{F})$ ist die Menge aller Erweiterungselemente von R .

Beweis. Die Erweiterungselemente bilden offenbar eine lin. M., da diese, wie wir bereits wissen, $\mathfrak{A}$, $\mathfrak{H} - \mathfrak{E}$, $\mathfrak{H} - \mathfrak{F}$ umfassen muß, ist $\mathfrak{A} + (\mathfrak{H} - \mathfrak{E}) + (\mathfrak{H} - \mathfrak{F})$ Teil von ihr; es bleibt zu zeigen, daß sie nicht größer ist.

Sei also f Erweiterungselement, nach dem obigen bedeutet das:

$$(f^*, \varphi - U\varphi) = (-if, \varphi + U\varphi), (f^* + if, \varphi) = (f^* - if, U\varphi)$$

(für alle φ von $\mathfrak{E}$, f^* fest). Sei U^{-1} die Inverse von U , mit $\mathfrak{F}$, $\mathfrak{E}$ als Definitionsbereich bzw. Wertevorrat. Es ist daher $P_{\mathfrak{F}}U = U$, $U^{-1}U = 1$ und $P_{\mathfrak{E}}U^{-1} = U^{-1}$, $UU^{-1} = 1$, soweit diese Operatoren Sinn haben, d. h. in $\mathfrak{E}$ bzw. $\mathfrak{F}$. Nunmehr ist

$$\begin{aligned} (f^* - if, U\varphi) &= (f^* - if, P_{\mathfrak{F}}U\varphi) = (P_{\mathfrak{F}}f^* - iP_{\mathfrak{F}}f, U\varphi) \\ &= (UU^{-1}P_{\mathfrak{F}}f^* - iUU^{-1}P_{\mathfrak{F}}f, U\varphi) = (U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f, \varphi), \end{aligned}$$

also

$$(\{f^* + if\} - \{U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f\}, \varphi) = 0.$$

Da dies für alle φ von $\mathfrak{E}$ gilt, so ist

$$\begin{aligned} P_{\mathfrak{E}}(\{f^* + if\} - \{U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f\}) &= 0, \\ \{P_{\mathfrak{E}}f^* + iP_{\mathfrak{E}}f\} - \{U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f\} &= 0, \\ -U^{-1}P_{\mathfrak{F}}f^* + P_{\mathfrak{E}}f^* &= -i(U^{-1}P_{\mathfrak{F}}f + P_{\mathfrak{E}}f), \end{aligned}$$

oder nach Anwendung von $-U$

$$P_{\mathfrak{F}}f^* - UP_{\mathfrak{E}}f^* = i(P_{\mathfrak{F}}f + UP_{\mathfrak{E}}f).$$

Weil allgemein $g - P_{\mathfrak{M}}g = P_{\mathfrak{F}-\mathfrak{M}}g = 0 \dots \text{mod } (\mathfrak{H} - \mathfrak{M})$ gilt, ist $P_{\mathfrak{E}}g = P_{\mathfrak{F}}g \dots \text{mod } ((\mathfrak{H} - \mathfrak{E}) \dot{+} (\mathfrak{H} - \mathfrak{F}))$. Darum folgt aus der obigen Gleichung

$$P_{\mathfrak{E}}f^* - UP_{\mathfrak{E}}f^* = i(P_{\mathfrak{E}}f + UP_{\mathfrak{E}}f) \dots \text{mod } ((\mathfrak{H} - \mathfrak{E}) \dot{+} (\mathfrak{H} - \mathfrak{F})).$$

Die linke Seite gehört ihrer Natur nach zu $\mathfrak{A}$, also gehört die rechte zu $\mathfrak{A} \dot{+} (\mathfrak{H} - \mathfrak{E}) \dot{+} (\mathfrak{H} - \mathfrak{F})$; ebenso gehört $P_{\mathfrak{E}}f - UP_{\mathfrak{E}}f$ seiner Natur nach zu $\mathfrak{A}$, aus beidem folgt, daß auch $P_{\mathfrak{E}}f$ zu $\mathfrak{A} \dot{+} (\mathfrak{H} - \mathfrak{E}) \dot{+} (\mathfrak{H} - \mathfrak{F})$ gehört. Da aber $f = P_{\mathfrak{E}}f \dots \bmod (\mathfrak{H} - \mathfrak{E})$ ist, gilt dasselbe von f — womit alles bewiesen ist.

Satz 30. $\mathfrak{A}$, $\mathfrak{H} - \mathfrak{E}$, $\mathfrak{H} - \mathfrak{F}$ sind lin. unabh.

Beweis. Wir müssen aus $f + \varphi + \psi = 0$, f von $\mathfrak{A}$, φ von $\mathfrak{H} - \mathfrak{E}$, ψ von $\mathfrak{H} - \mathfrak{F}$, folgern, daß $f = \varphi = \psi = 0$ ist, oder auch: daraus, daß φ zu $\mathfrak{H} - \mathfrak{E}$, ψ zu $\mathfrak{H} - \mathfrak{F}$, $\varphi + \psi$ zu $\mathfrak{A}$ gehört, $\varphi = \psi = 0$. Sei also diese letztere Prämisse erfüllt.

φ , ψ , $\varphi + \psi$ sind nach früher Gesagtem Erweiterungselemente, denen $R\,i\varphi$, $-i\psi$, $R(\varphi + \psi)$ (dieses hat Sinn!) zugeordnet. Da $\varphi + \psi$ offenbar auch $i(\varphi - \psi)$ zugeordnet wird, muß dieses $= R(\varphi + \psi)$ sein. Weil φ , ψ $i\varphi$ bzw. $-i\psi$ zugeordnet sind, so muß

$$(i\varphi, \varphi + \psi) = (\varphi, R(\varphi + \psi)) = (\varphi, i(\varphi - \psi)), \quad 2i(\varphi, \varphi) = 0, \quad \varphi = 0,$$

$$(-i\psi, \varphi + \psi) = (\psi, R(\varphi + \psi)) = (\psi, i(\varphi - \psi)), \quad -2i(\psi, \psi) = 0, \quad \psi = 0$$

gelten. —

Wir können also jedes Erweiterungselement f auf eine und nur eine Weise in $f = f^0 + f^+ + f^-$, f^0 von $\mathfrak{A}$, f^+ von $\mathfrak{H} - \mathfrak{E}$, f^- von $\mathfrak{H} - \mathfrak{F}$, zerlegen, und umgekehrt ist jedes solche f ein Erweiterungselement. Wir nennen f^0, f^+, f^- bzw. die 0-, +-, -Komponente von f . Da f^0, f^+, f^- $Rf^0, if^+, -if^-$ zugeordnet wird, so ordnet R dem $f = f^0 + f^+ + f^-$ $f^* = Rf^0 + if^+ - if^-$ zu. — Diese Zerlegung ist von Wichtigkeit, weil sie eine nähere Kennzeichnung der 0-, +-, -Klassen erlaubt:

Satz 31. Ein Erweiterungselement f gehört zur 0-, +-, -Klasse, je nachdem ob $|f^+| =, >, < |f^-|$ ist.

Beweis. Wegen $f = f^0 + f^+ + f^-$, $f^* = Rf^0 + if^+ - if^-$ und der Zugehörigkeit von if^+ , $-if^-$ zu f^+ bzw. f^- ist:

$$\begin{aligned} (f^*, f) &= (Rf^0 + if^+ - if^-, f^0 + f^+ + f^-) \\ &= (Rf^0, f^0) + \{(Rf^0, f^+) + (if^+, f^0)\} + \{(Rf^0, f^-) + (-if^-, f_0)\} \\ &\quad + \{(if^+, f^-) + (-if^-, f^+)\} + \{(if^+, f^+) + (-if^-, f^-)\} \\ &= (f^0, Rf^0) + \{(f^0, if^+) + (if^+, f^0)\} + \{(f^0, -if^-) + (-if^-, f_0)\} \\ &\quad + \{(f^+, -if^-) + (-if^-, f^+)\} + i\{|f^+|^2 - |f^-|^2\}. \end{aligned}$$

Das erste Glied und die drei ersten $\{ \}$ sind offenbar reell, die letzte $\{ \}$ offenbar rein imaginär, woraus

$$\Im(f^*, f) = |f^+|^2 - |f^-|^2$$

folgt — also die Behauptung.

VII. Die Defektindizes.

Definition 15. $R, U, \mathfrak{A}, \mathfrak{E}, \mathfrak{F}$ mögen die bisherige Bedeutung haben, wie z. B. in Satz 22. m, n seien die Dimensionszahlen von $\mathfrak{F} - \mathfrak{E}$ bzw. $\mathfrak{F} - \mathfrak{F}$.

Wir nennen m, n die Defektindizes von R . (Es ist $m, n = 0, 1, 2, \dots, \infty$).

Diese Zahlen werden bei den Fragen der Max. und Hypermax. eine entscheidende Rolle spielen. Zwischen ihnen und den drei Klassen der Erweiterungselemente besteht ein enger Zusammenhang, den wir nun dartun werden. Zunächst verallgemeinern wir den Begriff der Dimensionszahl in $\mathfrak{F}$.

Definition 16. $\mathfrak{M}$ sei eine (nicht notwendig abg.!) lin. M., $\mathfrak{U}$ eine Teilmenge von $\mathfrak{F}$ (die nicht einmal lin. M. zu sein braucht). Unter der Dimensionszahl von $\mathfrak{U} \dots \bmod \mathfrak{M}$ verstehen wir die obere Grenze aller (endlichen!) n , für die es n lin. unabh. Elemente gibt, die zusammen mit ihrer lin. M.⁴⁴) (mit eventueller Ausnahme der 0⁴⁵) ganz in $\mathfrak{U}$ liegen und außerhalb von $\mathfrak{M}$. n ist also eine der Zahlen $0, 1, 2, \dots, \infty$.

Übrigens wird im folgenden $\mathfrak{U}$ doch nicht ganz willkürlich sein, es wird vielmehr stets die folgenden Eigenschaften haben: Erstens ist es eine Restklassenmenge $\bmod \mathfrak{M}$, d. h. von zwei f, g , $f = g \dots \bmod \mathfrak{M}$, gehört keins oder beide zu $\mathfrak{U}$, und zweitens gehören mit f auch alle af ($a \neq 0$) zu $\mathfrak{U}$. Man sieht sofort, daß unter diesen $\mathfrak{U}$ die leere Menge und $\mathfrak{M}$ die einzigen sind, die die Dimensionszahl 0 haben. Weiter ist es klar, daß die Dimensionszahl einer abg. lin. M. $\mathfrak{U} \dots \bmod (0)$ mit der von Kap. VI übereinstimmt.

Die $+-$, $--$ und 0 -Klassen sind solche Mengen $\mathfrak{U} \bmod \mathfrak{A}$, ebenso die Vereinigungsmengen irgendwelcher von ihnen; da die zwei ersten 0 nicht enthalten, wohl aber die dritte, so müssen sie im Falle der 0-Dimensionalität leer bzw. gleich $\mathfrak{A}$ sein.

Satz 32. *Die $+$ -Klasse, sowie ihre Vereinigung mit der 0 -Klasse, ist $\bmod \mathfrak{A}$ m -dimensional; die $-$ -Klasse, sowie ihre Vereinigung mit der 0 -Klasse ist $\bmod \mathfrak{A}$ n -dimensional; die 0 -Klasse ist $\bmod \mathfrak{A}$ $\text{Min}(m, n)$ -dimensional.*

Beweis. Sei $p \leq m$ und endlich (es ist ja eventuell $m = \infty$), und $f_1, \dots, f_p$ lin. unabh. Elemente von $\mathfrak{F} - \mathfrak{E}$ (das ja m -dimensional ist). Die ganze lin. M. derselben liegt in $\mathfrak{F} - \mathfrak{E}$, also, mit Ausnahme der 0, in der $+$ -Klasse — und somit von selbst außerhalb $\mathfrak{A}$. Also hat die $+$ -Klasse $\geq m$ -Dimensionen $\bmod \mathfrak{A}$, wenn wir noch zeigen, daß ihre Vereinigung

⁴⁴) Die wegen der Endlichkeit von n von selbst abg. ist.

⁴⁵) Diese Ausnahme wird sich im folgenden als praktisch erweisen.

mit der 0-Klasse $\leq m$ hat, sind die zwei ersten Behauptungen bewiesen. Für $m = \infty$ ist dies trivial, bei endlichem m ist zu zeigen: wenn $f_1, \dots, f_{m+1}$ lin. unabh. Erweiterungselemente sind, so liegt gewiß ein

$$a_1 f_1 + \dots + a_{m+1} f_{m+1} \neq 0$$

in $\mathfrak{U}$ oder in der $-$ -Klasse.

Sei $f_\mu = f_\mu^0 + f_\mu^+ + f_\mu^-$ ($\mu = 1, \dots, m+1$), die f_μ^+ liegen alle in dem m -dimensionalen $\mathfrak{S} - \mathfrak{E}$, sind also nicht lin. unabh. Sei etwa

$$a_1 f_1^+ + \dots + a_{m+1} f_{m+1}^+ = 0$$

(ohne $a_1 = \dots = a_{m+1} = 0$, also $a_1 f_1 + \dots + a_{m+1} f_{m+1} \neq 0$!), dann ist $a_1 f + \dots + a_{m+1} f_{m+1} = g^0 + g^-$, g^0 von $\mathfrak{U}$, g^- von $\mathfrak{S} - \mathfrak{F}$. Für $g^- = 0$ gehört es also zu $\mathfrak{U}$, für $g^- \neq 0$ (nach Satz 31) zur $-$ -Klasse.

Damit sind die zwei ersten Behauptungen bewiesen, die zwei folgenden ergeben sich ebenso durch Vertauschen von $+$ und $-$. Es bleibt die letzte. Da die Dimensionszahl der 0-Klasse mod $\mathfrak{U}$ allenfalls $\leq m$ und $\leq n$ ist, genügt es zu zeigen, daß sie $\geq \text{Min.}(m, n)$ ist: d. h. zu jedem $p \leq \text{Min.}(m, n)$ solche lin. unabh. $f_1, \dots, f_p$ anzugeben, daß jedes $a_1 f_1 + \dots + a_p f_p \neq 0$ zur 0-Klasse gehört, aber nicht zu $\mathfrak{U}$.

In $\mathfrak{S} - \mathfrak{E}$ gibt es gewiß p lin. unabh. Elemente, die wir gleich „orthogonalisieren“ (vgl. Satz 8), d. h. norm. orth. annehmen können: $\varphi_1, \dots, \varphi_p$; ebenso seien $\psi_1, \dots, \psi_p$ norm. orth. aus $\mathfrak{S} - \mathfrak{F}$. Wir setzen

$$f_1 = \varphi_1 + \psi_1, \dots, f_p = \varphi_p + \psi_p,$$

dann ist (wenn nicht $a_1 = \dots = a_p = 0$)

$$a_1 f_1 + \dots + a_p f_p = (a_1 \varphi_1 + \dots + a_p \varphi_p) + (a_1 \psi_1 + \dots + a_p \psi_p)$$

(nach Satz 30) sicher nicht in $\mathfrak{U}$, also insbesondere $\neq 0$; und wegen

$$|a_1 \varphi_1 + \dots + a_p \varphi_p|^2 = |a_1 \psi_1 + \dots + a_p \psi_p|^2 = |a_1|^2 + \dots + |a_p|^2$$

gehört es (nach Satz 31) zur 0-Klasse. Die $f_1, \dots, f_p$ sind also lin. unabh. und leisten alles was wir wünschen.

Satz 33. *R ist dann und nur dann max., wenn $m = 0$ oder $n = 0$ ist, und hypermax., wenn $m = n = 0$ ist. Oder auch: wenn $\mathfrak{E} = \mathfrak{S}$ oder $\mathfrak{F} = \mathfrak{S}$ ist bzw. wenn $\mathfrak{E} = \mathfrak{F} = \mathfrak{S}$ ist.*

Beweis. Nach Satz 11 bedeutet die Max.: 0-Klasse gleich $\mathfrak{U}$, nach Definition 9 die Hypermax.: außerdem noch die $+-$ und die $-$ -Klasse leer. Nach der Bemerkung von Satz 32 bedeutet ersteres: 0-Klasse 0-dimensional, und letzteres: $+-$ und $-$ -Klasse 0-dimensional. Nach Satz 32 schließlich heißt dies: $\text{Min.}(m, n) = 0$, d. h. $m = 0$ oder $n = 0$ bzw. $m = 0$ und $n = 0$.

Die zweite Formulierung folgt sofort aus der ersten. —

Der Gegensatz von Max. und Hypermax. tritt nunmehr klar hervor, ebenso die Rolle von m, n . Zu Satz 32 sei noch bemerkt: Wenn wir R durch $aR + b1$ ersetzen (a, b reell, $a \neq 0$; dies ist wieder ein abg. lin. H. O. mit demselben Definitionsbereich wie R), so ändern sich $\mathfrak{E}, \mathfrak{F}$ auf unübersichtliche Weise, und wir sehen daher zunächst nicht, was mit m, n geschieht. Aber $\mathfrak{M}$, die Menge der Erweiterungselemente, sowie die $0-, +-, --$ -Klassen ändern sich bei $a > 0$ offenbar überhaupt nicht, und bei $a < 0$ werden nur die zwei letzteren vertauscht. Daher hat Satz 32 die Konsequenz, daß sich m, n auch nicht ändern bzw. nur vertauscht werden ⁴⁶⁾.

VIII. Der Fortsetzungsprozeß.

Bereits am Schlusse des Kap. VI hatten wir eine Übersicht über alle möglichen Forts. eines abg. lin. H. O. R gewonnen; wir können jetzt entscheiden, wann und wie eine Forts. zu max. bzw. hypermax. Operatoren gelingt.

Satz 34. *Ein abg. lin. H. O. R mit den Defektindizes m, n kann dann und nur dann zu einem mit den Defektindizes m', n' fortgesetzt werden, wenn es ein $p (= 0, 1, 2, \dots, \infty)$ mit $m' + p = m, n' + p = n$ gibt.*

Beweis. Wenn wir den Fortsetzungsprozeß nach Satz 27 betrachten ($U, \mathfrak{M}, \mathfrak{E}, \mathfrak{F}$ mögen den üblichen Sinn haben, $\mathfrak{M}, \mathfrak{N}$ wie in Satz 27), so sehen wir: $\mathfrak{S} - \mathfrak{E}, \mathfrak{S} - \mathfrak{F}$ sind m - bzw. n -dimensional, und die Frage ist, ob es zwei gleichviel- (etwa p -) dimensionale abg. lin. M. $\mathfrak{M}, \mathfrak{N}$ gibt, die

⁴⁶⁾ Sei $\bar{R}$ der folgende Operator: $\bar{R}f$ hat Sinn für alle Erweiterungselemente f von R , und zwar ist es dann dem zugeordneten f^* gleich. $\bar{R}$ ist offenbar abg. lin. und Forts. von R . (Wenn R Hypermax. ist, so ist offenbar $\bar{R} = R$, andernfalls ist es nicht einmal ein H. O., denn wegen $m > 0$ oder $n > 0$ ist dann die $+-$ oder die $--$ -Klasse nicht leer und in diesen ist $(\bar{R}f, f) = (f^*, f)$ nicht reell.) Nach Satz 28 sind $\mathfrak{S} - \mathfrak{E}$ und $\mathfrak{S} - \mathfrak{F}$ die Menge der Lösungen von $\bar{R}f = if$ bzw. $= -if$. Also ist für diese Gleichungen m bzw. n die Höchstzahl von lin. unabh. Lösungen.

Nach dem vorhin Gesagten ist aber die Höchstzahl der lin. unabh. Lösungen von $a\bar{R}f + bf = if$ gleich m bzw. n für $a >$ bzw. < 0 . Die obige Gleichung besagt: $\bar{R}f = \varrho f$, $\varrho = -\frac{b}{a} + \frac{1}{a}i$. Wir können also sagen: Wenn ϱ eine nicht reelle Zahl ist, so hat die Gleichung

$$\bar{R}f = \varrho f$$

die Höchstzahl m bzw. n von lin. unabh. Lösungen, wenn $\Im \varrho >$ bzw. < 0 ist. (Bei reellem ϱ ist die Höchstzahl bekanntlich die „Vielfachheit des Punkteigenwertes“ ϱ . $Rf = \varrho f$ muß bei nicht reellem ϱ unlösbar sein — außer durch $f = 0$; da sonst $(Rf, f) = \varrho (f, f)$ nicht reell ist.)

Diese Tatsachen stehen in Beziehung zu gewissen Resultaten von Carleman, vgl. das Zitat von Anm. 4).

Teile von denselben sind, und zwar so, daß $\mathfrak{S} - \mathfrak{E} - \mathfrak{M}$, $\mathfrak{S} - \mathfrak{F} - \mathfrak{N}$ m' - bzw. n' -dimensional ist.

Allgemeiner ist die Frage diese: Wann hat eine k -dimensionale abg. lin. M. $\mathfrak{R}$ ein p -dimensionales Teil (eine abg. lin. M.) $\mathfrak{P}$, so daß $\mathfrak{R} - \mathfrak{P}$ k' Dimensionen hat? Wir werden sehen: wenn $k = k' + p$ ist. Wenn wir hierin für $\mathfrak{R}$, $\mathfrak{P}$, k , k' $\mathfrak{S} - \mathfrak{E}$, $\mathfrak{M}$, m , m' und $\mathfrak{S} - \mathfrak{F}$, $\mathfrak{N}$, n , n' einsetzen, haben wir unseren Satz.

Notwendig ist die Bedingung: denn seien $\mathfrak{P}$, $\mathfrak{R} - \mathfrak{P}$ p - bzw. k' -dimensional, dann werden sie (als abg. lin. M.) von den norm. orth. Systemen $\varphi_1, \dots, \varphi_p$ bzw. $\psi_1, \dots, \psi_{k'}$ aufgespannt, also $\mathfrak{R} = \mathfrak{P} + (\mathfrak{R} - \mathfrak{P})$ von $\varphi_1, \dots, \varphi_p, \psi_1, \dots, \psi_{k'}$ (diese sind orth., weil es $\mathfrak{P}$, $\mathfrak{R} - \mathfrak{P}$ sind), d. h. es ist $p + k'$ -dimensional, also $k = k' + p$. (Hierbei kann auch k' oder $p = \infty$ sein!) — Hinreichend ist sie ebenfalls, denn sei $\mathfrak{R}$ $k = k' + p$ -dimensional, dann wird es von einem norm. orth. System $\varphi_1, \dots, \varphi_p, \psi_1, \dots, \psi_{k'}$ aufgespannt (wieder darf k' oder $p = \infty$ sein!). Die von $\varphi_1, \dots, \varphi_p$ aufgespannte abg. lin. M. sei $\mathfrak{P}$, dann wird $\mathfrak{R} - \mathfrak{P}$ von $\psi_1, \dots, \psi_{k'}$ aufgespannt, und wir sind am Ziele.

Satz 35. *R sei ein abg. lin. H.O. mit den Defektindizes m, n . Zu einem max. H.O. kann R stets fortgesetzt werden. Wenn $m \neq n$ ist, so gibt es unter seinen Forts. keinen hypermax.; wenn $m = n$ ist, so ist, falls diese endlich sind, jede max. Forts. hypermax., wenn sie $= \infty$ sind, so gibt es unter den max. Forts. sowohl solche, die hypermax. sind, als auch solche, die es nicht sind.*

Beweis. Nach Satz 33 bedeutet Max.: mindestens ein Defektindex $= 0$, Hypermax.: beide $= 0$, Max. aber nicht Hypermax.: genau einer $= 0$. Wenn man hierauf Satz 34 anwendet, gewinnt man genau die Behauptungen unseres Satzes.

Zusatz. *Wenn R nicht schon max. ist, so ist die max. Forts. (und, wenn es solche gibt, die hypermax. sowie die max. aber nicht hypermax. Forts.) auf Kontinuum viele verschiedene Arten möglich.*

Beweis. Unsere Forts.-Methoden kamen (nach Satz 27) auf das Erweitern der $\mathfrak{S} - \mathfrak{E}$, $\mathfrak{S} - \mathfrak{F}$ um zwei abg. lin. M. $\mathfrak{M}$, $\mathfrak{N}$ heraus, die bei nicht max. $R > 0$ Dimensionen hatten. Aber dabei spielte noch eine längentreue Abbildung U von $\mathfrak{M}$ auf $\mathfrak{N}$ eine Rolle, und diese kann auf Kontinuum viele Arten gewählt werden: zu jedem $\mathfrak{N}$ aufspannenden norm. orth. System $\psi_1, \dots, \psi_p$ gehörte (nach Satz 26) ein U , und derer gibt es Kontinuum viele (weil z. B. ψ_1 durch jedes $\alpha \psi_1$, $|\alpha| = 1$, ersetzt werden kann)⁴⁷⁾.

⁴⁷⁾ Mehr als Kontinuum viele Forts. kann es nicht geben, da es nur Kontinuum viele abg. Operatoren in $\mathfrak{S}$ gibt — wie man aus der Separabilität von $\mathfrak{S}$ leicht folgert.

Wenn R nur H. O. ist, so sind seine max. und hypermax. Forts., wie wir wissen, dieselben wie beim abg. lin. H. O. $\tilde{R}$, auf welchen die soeben erzielten Resultate anwendbar sind. Hat insbesondere R nur eine max. Forts., so muß schon $\tilde{R}$ max. sein, ebenso bei den hypermax. Forts.

Wir wenden uns nunmehr der näheren Untersuchung der hypermax. H. O. zu (Kap. IX), und dann der max., aber nicht hypermax. H. O. (Kap. X).

IX. Die hypermaximalen Operatoren und die Eigenwertdarstellung.

Wenn R hypermaximal ist, so ist $\mathfrak{E} = \mathfrak{F} = \mathfrak{S}$, d. h. die Cayleysche Transformierte U ist unitär: sie ist überall sinnvoll und hat eine (ebenfalls unitäre) überall sinnvolle Inverse U^{-1} . Die hypermax. H. O. R entsprechen als ein-eindeutig den unitären Operatoren U , die die Bedingung von Satz 24 erfüllen, d. h. für die die Menge der $\varphi - U\varphi$ überall dicht ist.

Schon beim Beweise von Satz 24 konstatierten wir: wenn die $\varphi - U\varphi$ überall dicht liegen, so folgt aus $U\varphi = \varphi$ $\varphi = 0$. Es gilt aber hier auch die Umkehrung. Wenn nämlich die $\varphi - U\varphi$ nicht überall dicht liegen, so spannen sie eine abg. lin. M. $\mathfrak{M} \neq \mathfrak{S}$ auf (sie selber bilden schon eine lin. Mann.!) — dann ist $\mathfrak{S} - \mathfrak{M} \neq (0)$. Sei also $f \neq 0$ ein Element von $\mathfrak{S} - \mathfrak{M}$. Alle $\varphi - U\varphi$ gehören zu $\mathfrak{M}$, sind also zu f orth.:

$$(f, \varphi - U\varphi) = (f, \varphi) - (f, U\varphi) = (Uf, U\varphi) - (f, U\varphi) = (Uf - f, U\varphi)$$

muß für alle φ verschwinden; daher ist $Uf - f = 0$, $Uf = f$.

Die in Frage kommenden unitären U können also auch so charakterisiert werden: für sie hat $Uf = f$ keine andere Lösung als die 0 (1 ist nicht „Punkteigenwert“!). Wir beherrschen somit die hypermax. H. O. R , wenn wir diese U kennen. Dies ist aber ein Problem für beschränkte Operatoren (U ist längentreu, also stetig), das auch mit den Methoden der Hilbertschen Theorie der beschränkten Formen (mit einem, auch jenem Ideenkreise angehörenden, Zusatz) bewältigt werden kann. Wir führen die notwendigen Überlegungen (in teilweiser Anlehnung an die einfache Methode von F. Riesz) im Anhang II durch und geben hier nur das Resultat an.

Definition 17. $a < x < b$ sei ein offenes Intervall (im folgenden wird a, b einmal 0, 1 und einmal $-\infty, \infty$ sein). Unter einer Zerlegung der Einheit (kurz: Z. d. E.) verstehen wir eine Schar von P. O. $E(\lambda)$ (λ läuft durch ganz a, b) mit den folgenden Eigenschaften:

α) Wenn $\lambda \leq \mu$ ist, so ist $E(\lambda)$ Teil von $E(\mu)$.

β) Wenn $\lambda \geq \lambda_0$, $\lambda \rightarrow \lambda_0$, so gilt für jedes f $E(\lambda)f \rightarrow E(\lambda_0)f$.

γ) Wenn $\lambda \rightarrow a$ oder $\rightarrow b$, so ist $E(\lambda)f \rightarrow 0$ bzw. $\rightarrow f$.

Wenn nun a, b das Intervall $0, 1$ ist und $E(\varrho)$ eine Z. d. E., so gibt es zu jedem f genau ein f^* , so daß für alle g

$$(f^*, g) = \int_0^1 e^{2\pi i \varrho} d(E(\varrho) f, g) \quad 48)$$

gilt. Der durch $Uf = f^*$ definierte Operator U ist unitär, und $Uf = f$ hat keine Lösung außer 0; und umgekehrt wird jedes U mit diesen Eigenschaften durch genau eine Z. d. E. $E(\lambda)$ auf die soeben skizzierte Weise erzeugt (vgl., wie gesagt, Anhang II).

Wir beweisen nun:

Satz 36. a, b sei das Intervall $-\infty, \infty$ und $F(\lambda)$ eine Z. d. E. Wenn das Integral

$$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda) f|^2$$

endlich ist⁴⁹⁾, so gibt es genau ein f^* , so daß für alle g

$$(f^*, g) = \int_{-\infty}^{\infty} \lambda d(F(\lambda) f, g)$$

(das Integral rechts ist absolut konvergent) gilt. Wir definieren einen Operator R für diese f , und zwar durch $Rf = f^*$.

Dieses R ist ein hypermax. H. O., und jeder hypermax. H. O. R kann mit Hilfe genau einer Z. d. E. auf diese Weise erzeugt werden.

Beweis. Den Z. d. E. $E(\varrho)$ fürs Intervall $0, 1$ werden die Z. d. E. $F(\lambda)$ fürs Intervall $-\infty, \infty$ durch die Zuordnung

$$E(\varrho) \Rightarrow F(\lambda) = E\left(-\frac{1}{\pi} \operatorname{arc} \operatorname{ctg} \lambda\right), \quad F(\lambda) \Rightarrow E(\varrho) = F(-\operatorname{ctg} \pi \varrho) \quad 50)$$

ein-eindeutig zugeordnet. Es genügt nun zu zeigen: Wenn $E(\lambda)$ die Z. d. E. eines unitären Operators U ist (vgl. die Ausführungen vor unserem Satze), so erzeugt $F(\lambda)$ nach der Anweisung des Satzes wirklich einen Operator R , und dieser ist die Cayleysche Transformierte von U . Nach dem, was wir über das Verhältnis der Cayleyschen Transformaten wissen, und über hypermax. H. O. sowie unitäre Operatoren und Z. d. E. vor dem Satze

⁴⁸⁾ Stieltjessches Integral, es ist absolut konvergent. Wir schreiben ϱ statt λ .

⁴⁹⁾ Mit λ wächst $F(\lambda)$, also auch $|F(\lambda) f|^2$, außerdem ist λ^2 nie negativ, daher ist dieses Integral seiner Natur nach entweder konvergent (endlich) oder eigentlich divergent $(+\infty)$.

⁵⁰⁾ Durch $\varrho = -\frac{1}{\pi} \operatorname{arc} \operatorname{ctg} \lambda$, $\lambda = -\operatorname{ctg} \pi \varrho$ werden die Intervalle $0 < \varrho < 1$ und $-\infty < \lambda < \infty$ ein-eindeutig aufeinander bezogen. (Beim $\operatorname{arc} \operatorname{ctg}$ ist immer der Wert $> -\pi$, < 0 zu nehmen.)

formuliert haben, folgt hieraus tatsächlich die Richtigkeit aller Behauptungen.

Sei also $F(\lambda)$ eine Z. d. E., $E(\varrho) = F(-\operatorname{ctg} \pi \varrho)$, U der dazugehörige unitäre Operator. Zunächst zeigen wir: Wenn $\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2 = C^2$ endlich ist, so kann das f^* unseres Satzes wirklich erzeugt werden.

Es sei $\lambda < \lambda' < \mu$. Da $E(\lambda)$ Teil von $E(\mu)$ ist, ist $E(\mu) - E(\lambda)$ ein P. O. und deshalb

$$\begin{aligned} |(\{E(\mu) - E(\lambda)\}f, g)| &= |(\{E(\mu) - E(\lambda)\}f, \{E(\mu) - E(\lambda)\}g)| \\ &\leq |\{E(\mu) - E(\lambda)\}f| \cdot |\{E(\mu) - E(\lambda)\}g|, \end{aligned}$$

und weil allgemein

$$\begin{aligned} |\{E(\mu) - E(\lambda)\}h|^2 &= (h, \{E(\mu) - E(\lambda)\}h) = (h, E(\mu)h) - (h, E(\lambda)h) \\ &= |E(\mu)h|^2 - |E(\lambda)h|^2 \end{aligned}$$

gilt,

$$\begin{aligned} |(\{E(\mu) - E(\lambda)\}f, g)| &\leq \sqrt{|E(\mu)f|^2 - |E(\lambda)f|^2} \sqrt{|E(\mu)g|^2 - |E(\lambda)g|^2}, \\ &|\lambda' \{E(\mu)f, g\} - (E(\lambda)f, g)| \\ &\leq \sqrt{\lambda'^2 \{|E(\mu)f|^2 - |E(\lambda)f|^2\}} \sqrt{|E(\mu)g|^2 - |E(\lambda)g|^2}. \end{aligned}$$

Da nun die beiden Integrale $\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2 = C^2$ und $\int_{-\infty}^{\infty} d|E(\lambda)g|^2 = |E(\lambda)g|^2 \Big|_{\lambda=-\infty}^{\lambda=\infty} = |g|^2$ absolut konvergieren, hat die bekannte Ungleichung

$$\begin{aligned} \sqrt{a_1 b_1} + \dots + \sqrt{a_p b_p} &\leq \sqrt{(a_1 + \dots + a_p)(b_1 + \dots + b_p)} \\ &(\text{alle } a_1, \dots, a_p, b_1, \dots, b_p \geq 0) \end{aligned}$$

erstens die Konsequenz, daß auch $\int_{-\infty}^{\infty} \lambda d(E(\lambda)f, g)$ absolut konvergiert, und zweitens daß

$$\left| \int_{-\infty}^{\infty} \lambda d(E(\lambda), g) \right| \leq \sqrt{\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2 \cdot \int_{-\infty}^{\infty} d|E(\lambda)g|^2} = C \cdot |g|$$

ist⁵¹⁾. Wir nennen dieses Integral für einen Augenblick $L(g)$ (f fest, g variabel), dann hat $L(g)$ die Eigenschaften:

$$L(a_1 g_1 + \dots + a_n g_n) = \bar{a}_1 L(g_1) + \dots + \bar{a}_n L(g_n), \quad |L(g)| \leq C \cdot |g|$$

(es ist linear und stetig). Nach einem Satze von F. Riesz gibt es darum

⁵¹⁾ Es genügt, sich die Definition des Stieltjesschen Integrals zu vergegenwärtigen. Die genannte Ungleichung ist nur eine veränderte Schreibweise der Schwarzschen.

genau ein f^* , so daß stets $(f^*, g) = L(g)$ ist⁵²⁾ — d. i. gerade unsere auf f^* bezügliche Behauptung.

R ist also wirklich herstellbar. Sei R^* die Cayleysche Transformierte von U , wir müssen noch $R = R^*$ beweisen. Es genügt aber zu zeigen, daß R Forts. von R^* ist. Denn wenn Rf, Rg Sinn haben, so geht (nach Definition in Satz 36) beim Vertauschen von f, g (Rf, g) in seine komplexkonjugierte über — also ist $(Rf, g) = (f, Rg)$. Daher ist R ein H. O. (sein Definitionsbereich enthält nach Annahme den von R^* , ist also überall dicht) und Forts. des (wegen der Unitarität von U) max. R^* , also $R = R^*$.

Sei also R^*f sinnvoll, wir müssen beweisen: Rf ist sinnvoll und $= R^*f$; oder was dasselbe ist: $\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2$ endlich und für alle g $(R^*f, g) = (Rf, g) = \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g)$. Daß Rf sinnvoll ist, bedeutet übrigens $f = \varphi - U\varphi$, dieses sei unser Ausgangspunkt.

$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2$, $\int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g)$ können wir, nach der Transformation $\lambda = -\operatorname{ctg} \pi \varrho$, auch so schreiben: $\int_0^1 \operatorname{ctg}^2 \pi \varrho d|E(\varrho)f|^2$, $\int_0^1 -\operatorname{ctg} \pi \varrho d(E(\varrho)f, g)$. Um hier $f = \varphi - U\varphi$ einsetzen zu können, müssen wir zunächst einige Integrale ausrechnen.

$$\begin{aligned} (\varphi \pm U\varphi, g) &= \int_0^1 d(E(\varrho)\varphi, g) \pm \int_0^1 e^{2\pi i \cdot \varrho} d(E(\varrho)\varphi, g) \\ &= \int_0^1 (1 \pm e^{2\pi i \cdot \varrho}) d(E(\varrho)\varphi, g), \end{aligned}$$

⁵²⁾ Man beweist diesen Satz so:

Da alle (f^*, g) vorgeschrieben sind, kann es nicht zwei verschiedene f^* geben, es genügt also die Existenz von einem zu beweisen. Sei $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System, $L(\varphi_n) = a_n$ ($n = 1, 2, \dots$). Es ist $g = \sum_{n=1}^{\infty} (g, \varphi_n) \cdot \varphi_n$ (Satz 7, β), also wegen der Linearität und Stetigkeit von $L(g)$ $L(g) = \sum_{n=1}^{\infty} (\overline{g, \varphi_n}) \cdot a_n$ (die Reihe muß konvergieren). Wäre nun $\sum_{n=1}^{\infty} |a_n|^2$ endlich, so würde $f^* = \sum_{n=1}^{\infty} a_n \varphi_n$ konvergieren (Satz 5), $(f^*, \varphi_n) = a_n$, und daraus folgt $L(g) = \sum_{n=1}^{\infty} (f^*, \varphi_n) \cdot (\overline{g, \varphi_n}) = (f^*, g)$ (Satz 7 γ): d. i. die Behauptung.

Nun ist $|L(g)| \leq C \cdot |g|$, also für $g = x_1 \varphi_1 + \dots + x_m \varphi_m$

$$\left| \sum_{n=1}^m a_n \bar{x}_n \right| \leq C \cdot \sqrt{\sum_{n=1}^m |x_n|^2},$$

und wenn wir $x_n = a_n$ ($n = 1, \dots, m$) setzen: $\sum_{n=1}^m |a_n|^2 \leq C^2$. Daher ist $\sum_{n=1}^{\infty} |a_n|^2$ endlich (und zwar $\leq C^2$).

$$\begin{aligned}
(\varphi - U\varphi, E(\varrho)g) &= \int_0^1 (1 - e^{2\pi i \cdot \varrho'} d(E(\varrho'))\varphi, E(\varrho)g) \\
&= \int_0^1 (1 - e^{2\pi i \cdot \varrho'}) d(E(\varrho)E(\varrho')\varphi, g) \\
&= \int_0^1 (1 - e^{2\pi i \cdot \varrho'}) d(E(\text{Min}(\varrho, \varrho'))\varphi, g) \quad {}^{53)} \\
&= \int_0^{\varrho} (1 - e^{2\pi i \cdot \varrho'}) d(E(\varrho')\varphi, g), \\
|E(\varrho)(\varphi - U\varphi)|^2 &= (\varphi - U\varphi, E(\varrho)(\varphi - U\varphi)) \\
&= \int_0^{\varrho} (1 - e^{2\pi i \cdot \varrho'}) d(E(\varrho')\varphi, \varphi - U\varphi) \quad {}^{54)} \\
&= \int_0^{\varrho} (1 - e^{2\pi i \cdot \varrho'}) d\left[\int_0^{\varrho'} (1 - e^{-2\pi i \cdot \varrho''}) d(E(\varrho'')\varphi, \varphi)\right] \quad {}^{55)} \\
&= \int_0^{\varrho} (1 - e^{2\pi i \cdot \varrho'})(1 - e^{-2\pi i \cdot \varrho'}) d(E(\varrho')\varphi, \varphi) \\
&= \int_0^{\varrho} 4 \sin^2 \pi \varrho' d|E(\varrho')\varphi|^2,
\end{aligned}$$

$$\begin{aligned}
\int_0^1 \text{ctg}^2 \pi \varrho d|E(\varrho)(\varphi - U\varphi)|^2 &= \int_0^1 \text{ctg}^2 \pi \varrho d\left[\int_0^{\varrho} 4 \sin^2 \pi \varrho' d|E(\varrho')\varphi|^2\right] \quad {}^{55)} \\
&= \int_0^1 \text{ctg}^2 \pi \varrho \cdot 4 \sin^2 \pi \varrho d|E(\varrho)\varphi|^2 = \int_0^1 4 \cos^2 \pi \varrho d|E(\varrho)\varphi|^2,
\end{aligned}$$

und dies ist endlich, weil es durch das endliche $\int_0^1 4 d|E(\varrho)\varphi|^2 = 4|\varphi|^2$ majorisiert wird. Weiter ist

$$\begin{aligned}
-\text{ctg} \pi \varrho d(E(\varrho)(\varphi - U\varphi), g) &= \int_0^1 -\text{ctg} \pi \varrho d\left[\int_0^{\varrho} (1 - e^{2\pi i \cdot \varrho'}) d(E(\varrho')\varphi, g)\right] \\
&= \int_0^1 -\text{ctg} \pi \varrho \cdot (1 - e^{2\pi i \cdot \varrho}) d(E(\varrho)\varphi, g) = \int_0^1 i(1 + e^{2\pi i \cdot \varrho}) d(E(\varrho)\varphi, g) \\
&= (i(\varphi + U\varphi), g) = (R^*(\varphi - U\varphi), g).
\end{aligned}$$

Damit ist alles Erforderliche bewiesen. —

⁵³⁾ Für $\varrho' > \varrho$ ist ja $(E(\text{Min}(\varrho, \varrho'))\varphi, g)$ konstant, also das Integral 0!

⁵⁴⁾ Man nehme mit $g = \varphi$ das komplex-Konjugierte der letzten Formel.

⁵⁵⁾ Nach einer allgemeinen Formel für Stieltjessche Integrale ist

$$\int_A^B p(\varrho) \cdot d\left[\int_A^{\varrho} q(\varrho') \cdot d\tau(\varrho')\right] = \int_A^B p(\varrho) \cdot q(\varrho) \cdot d\tau(\varrho).$$

Wir haben für hypermax. H. O. die der Hilbertschen Spektral-darstellung oder Eigenwertdarstellung (bei beschränkten Formen, vgl. das Zitat von Anm. 7)) analoge Normalform gewonnen, dabei aber für diese bereits alle Eigenwertdarstellungen verbraucht, so daß es für max. und nicht hypermax. H. O. gewiß keine solchen geben kann. (Allerdings haben wir noch kein Beispiel eines derartigen H. O. konstruiert.) Darum sehen wir von der kaum neue Gesichtspunkte bietenden weiteren Untersuchung der hypermax. H. O. ab und wenden uns den max., aber nicht hypermax. H. O. zu.

X. Maximale, aber nicht hypermaximale Operatoren.

Ein solcher H. O. R (m, n seine Defektindizes) ist durch $m = 0$, $n > 0$ oder $m > 0$, $n = 0$ charakterisiert; da aber beim Ersetzen von R durch $-R$ m, n vertauscht werden, genügt es $m = 0$, $n > 0$ zu betrachten. Wir bilden $U, \mathfrak{E}, \mathfrak{F}$, es ist $\mathfrak{E} = \mathfrak{S}$ und $\mathfrak{G}_0 = \mathfrak{S} - \mathfrak{F}$ ist n -dimensional, U bildet $\mathfrak{S}$ längentreu auf $\mathfrak{F}$ ab.

Das durch U^ν ($\nu = 0, 1, 2, \dots$) vermittelte Bild von $\mathfrak{G}_0$ sei $\mathfrak{G}_\nu$ (auch eine abg. lin. M.). Für $\nu > 0$ ist $\mathfrak{G}_\nu$ Teil von $\mathfrak{F}$, also orth. zu $\mathfrak{G}_0$; wenn wir hierauf die Abbildung U^k ($k = 0, 1, 2, \dots$) anwenden, so werden die Bilder $\mathfrak{G}_{k+\nu}, \mathfrak{G}_k$ zueinander orth. (weil (f, g) dabei invariant ist, vgl. Satz 10). Also sind allgemein $\mathfrak{G}_k, \mathfrak{G}_l$ ($k, l = 0, 1, 2, \dots$) für $k \neq l$ orth. Die $\mathfrak{G}_0, \mathfrak{G}_1, \mathfrak{G}_2, \dots$ mögen die abg. lin. M. $\mathfrak{M}$ aufspannen, wir setzen $\mathfrak{S} - \mathfrak{M} = \mathfrak{N}$.

Daß U $\mathfrak{G}_\nu$ auf $\mathfrak{G}_{\nu+1}$ abbildet (natürlich ein-eindeutig, längentreu) wissen wir schon, wir wollen nun zeigen, daß $\mathfrak{N}$ (ebenso) auf sich selbst abgebildet wird. $\mathfrak{N}$ ist, wie man leicht einsieht, die Menge aller zu allen $\mathfrak{G}_\nu$ orth. Elemente von $\mathfrak{S}$, sein Bild also (wegen der Invarianz von (f, g)) die Menge aller zu allen Bildern der $\mathfrak{G}_\nu$ orth. Elemente des Bildes von $\mathfrak{S}$: dies sind die zu allen $\mathfrak{G}_{\nu+1}$ orth. Elemente von $\mathfrak{F} = \mathfrak{S} - \mathfrak{G}_0$, also die zu $\mathfrak{G}_0$ und allen $\mathfrak{G}_{\nu+1}$ orth. Elemente von $\mathfrak{S}$. Also wirklich wieder $\mathfrak{N}$.

Sei nun $\varphi_1^0, \dots, \varphi_p^0$ ein die abg. lin. M. $\mathfrak{G}_0$ aufspannendes norm. orth. System (eventuell $p = \infty!$), und $U^\nu \varphi_q^0 = \varphi_q^\nu$ ($\nu = 0, 1, 2, \dots, q = 1, \dots, p$), dann bilden die $\varphi_1^\nu, \dots, \varphi_p^\nu$ auch ein norm. orth. System und spannen $\mathfrak{G}_\nu$ auf. Für $k \neq l$ liegt φ_q^k in $\mathfrak{G}_k$, $\varphi_{q'}^l$ in $\mathfrak{G}_l$, also sind beide orth., d. h. alle φ_q^ν ($\nu = 0, 1, 2, \dots, q = 1, \dots, p$) bilden auch ein norm. orth. System — und zwar spannen sie dasselbe auf, wie die $\mathfrak{G}_\nu : \mathfrak{M}$.

$\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$ mögen $\mathfrak{M}_q$ aufspannen, $\varphi_q^1, \varphi_q^2, \dots, \overline{\mathfrak{M}}_q$ (alles als abg. lin. M., $q = 1, \dots, p$), dann spannen $\mathfrak{M}_1, \dots, \mathfrak{M}_p$ wieder $\mathfrak{M}$ auf, und $\overline{\mathfrak{M}}_q$ ist Teil von $\mathfrak{M}_p$. Da $U\varphi_q^\nu = \varphi_q^{\nu+1}$ ist, bildet U $\mathfrak{M}_q$ auf $\overline{\mathfrak{M}}_q$ ab. Zusammenfassend haben wir also:

$\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ spannen $\mathfrak{H}$ auf (wie $\mathfrak{M}, \mathfrak{N}$), U bildet sie bzw. auf $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ ab, d. h. auf Teilmengen von sich selbst. Ferner sind die $\mathfrak{M}_1, \dots, \mathfrak{M}_p$ (als Teile von $\mathfrak{M}$) zu $\mathfrak{N}$ orth., und auch zueinander (weil es die $\varphi_q^k, \varphi_{q''}^l$ mit $q' \neq q''$ sind). Hieraus folgt, daß jedes $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ das U reduziert, also (Satz 25) auch die Cayleysche Transformierte R . —

Wir betrachten die in $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ liegenden Teile von R und U : es sind bzw. $R_1, \dots, R_p, S$ und bzw. $U_1, \dots, U_p, V$; R ist mit Hilfe von $R_1, \dots, R_p, S$ nach Satz 21 zu charakterisieren. Nun sind die abg. lin. M. $\mathfrak{M}_1, \dots, \mathfrak{M}_p$ alle unendlichvieldimensional ($\mathfrak{M}_q$ wird ja vom norm. orth. System $\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$ aufgespannt), also lauter Hilbertsche Räume (d. h. sie genügen, bei unveränderter Definition von $af, f+g, (f, g)$ den Bedingungen A bis E von Kap. I ⁵⁶⁾), daher dürfen wir in ihnen alle für $\mathfrak{H}$ ausgebildeten Begriffe ohne weiteres verwenden. Also ist R_q in $\mathfrak{M}_q$ ein H. O. (insbesondere ist sein Definitionsbereich in $\mathfrak{M}_q$ überall dicht, denn er besteht aus den Projektionen des in $\mathfrak{H}$ überall dichten Definitionsbereiches von R), U_q ist längentreu und die Cayleysche Transformierte von R_q . In $\mathfrak{N}$ sind drei Fälle möglich: erstens kann es ∞ -dimensional, also ein Hilbertscher Raum sein, dann ist wieder S ein H. O., V unitär (es bildet ja $\mathfrak{N}$ auf $\mathfrak{N}$ ab) und seine Cayleysche Transformierte — also S hypermax.; zweitens kann $\mathfrak{N}$ N ($=1, 2, \dots$)-dimensional, also ein N -dimensionaler (komplexer) Euklidischer Raum sein (vgl. Anm. ⁵⁶⁾), dann ist S ein H. O. in ihm — d. h. ein gewöhnlicher, endlichvieldimensionaler Hermitescher Operator; drittens kann $\mathfrak{N}$ 0-dimensional sein, d. h. $= (0)$ — dann brauchen wir es gar nicht zu berücksichtigen.

Bei $\mathfrak{N}$ und S geschieht also sicher nichts Neues, wir wenden uns darum den $\mathfrak{M}_q, R_q$ und U_q ($q=1, \dots, p$) zu. Diese benehmen sich alle gleich: R_q ist Cayleysche Transformierte von U_q , und es gibt ein in $\mathfrak{M}_q$ vollst. (d. h. $\mathfrak{M}_q$ aufspannendes) norm. orth. System $\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$, so daß $U_q \varphi_q^r = U \varphi_q^r = \varphi_q^{r+1}$ ist, also $U_q \left(\sum_{r=0}^{\infty} x_r \varphi_q^r \right) = \sum_{r=0}^{\infty} x_r \varphi_q^{r+1}$ ($\sum_{r=0}^{\infty} |x_r|^2$ endlich); hierdurch ist aber U_q und somit R_q völlig festgelegt (wenn man $\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$ als bekannt ansieht). Indem wir die Isomorphie von $\mathfrak{M}_q$ mit $\mathfrak{H}$ ausnützen, können wir auch sagen, daß R_q mit dem folgenden Operator $\bar{R}$ in $\mathfrak{H}$ isomorph sein muß: $\bar{R}$ ist die Cayleysche Transformierte von $\bar{U}$, und $\bar{U}$ ist

⁵⁶⁾ A bis C und E genügt offenbar jede abg. lin. M., D aber ist der Unendlichvieldimensionalität gleichwertig. Eine endlichvieldimensionale abg. lin. M. besteht, falls sie 0 Dimensionen hat, aus der 0 allein; wenn sie $N=1, 2, \dots$ Dimensionen hat, so wird sie durch ein norm. orth. System $\varphi_1, \dots, \varphi_N$ aufgespannt, sie besteht also aus den $x_1 \varphi_1 + \dots + x_N \varphi_N$, d. h. sie ist der N -dimensionale (komplexe) Euklidische Raum aller $x_1, \dots, x_N$.

(wenn $\psi_0, \psi_1, \psi_2, \dots$ ein festes vollst. norm. orth. System in $\mathfrak{H}$ ist) durch $\bar{U} \left(\sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu} \right) = \sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu+1}$ ($\sum_{\nu=0}^{\infty} |x_{\nu}|^2$ endlich) definiert. —

Dabei wäre aber zuerst noch zu zeigen, daß $\bar{R}$ wirklich existiert, d. h. daß das offenbar längentreue $\bar{U}$ der Bedingung von Satz 24 genügt. Wenn das der Fall ist, so ist $\bar{R}$, da $\bar{U}$ in ganz $\mathfrak{H}$ definiert (also $\mathfrak{E} = \mathfrak{H}$, $m = 0$) ist, ein max. H. O. Der Wertevorrat von $\bar{U}$, $\mathfrak{F}$ ist die Menge aller $\sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu+1}$, d. h. die von den $\psi_1, \psi_2, \dots$ aufgespannte abg. lin. M. — also ist $\mathfrak{H} - \mathfrak{F}$ die Menge aller $a \psi_1$, also 1-dimensional. Folglich hat $\bar{R}$ die Defektindizes 0, 1: es ist das erste Beispiel eines max., aber nicht hypermax. H. O.

Zunächst ist aber noch seine Existenz zu beweisen, d. h. die überall Dichtigkeit der Menge aller $\varphi - \bar{U} \varphi$. Da sie eine lin. M. ist, genügt es zu zeigen, daß alle Elemente eines vollst. norm. orth. Systems Häufungspunkte von ihr sind — etwa alle ψ_{ν} . Nun ist

$$\begin{aligned} & \left(\psi_{\nu} + \frac{l-1}{l} \psi_{\nu+1} + \dots + \frac{1}{l} \psi_{\nu+l-1} \right) - \bar{U} \left(\psi_{\nu} + \frac{l-1}{l} \psi_{\nu+1} + \dots + \frac{1}{l} \psi_{\nu+l-1} \right) \\ &= \left(\psi_{\nu} + \frac{l-1}{l} \psi_{\nu+1} + \dots + \frac{1}{l} \psi_{\nu+l-1} \right) - \left(\psi_{\nu+1} + \frac{l-1}{l} \psi_{\nu+2} + \dots + \frac{1}{l} \psi_{\nu+l} \right) \\ &= \psi_{\nu} - \frac{1}{l} (\psi_{\nu+1} + \dots + \psi_{\nu+l}), \end{aligned}$$

und dies unterscheidet sich von ψ_{ν} um $\frac{1}{l} (\psi_{\nu+1} + \dots + \psi_{\nu+l})$, wobei

$$\left| \frac{1}{l} (\psi_{\nu+1} + \dots + \psi_{\nu+l}) \right|^2 = \frac{1}{l}, \quad \left| \frac{1}{l} (\psi_{\nu+1} + \dots + \psi_{\nu+l}) \right| = \frac{1}{\sqrt{l}}$$

ist, also (bei geeignetem l) beliebig klein. Damit ist alles bewiesen.

Wir formulieren:

Satz 37. Sei $\psi_0, \psi_1, \psi_2, \dots$ ein vollst. norm. orth. System. Es gibt dann genau einen H. O. $\bar{R}$, dessen Cayleysche Transformierte $\bar{U}$ durch

$$\bar{U} \left(\sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu} \right) = \sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu+1} \quad \left(\sum_{\nu=0}^{\infty} |x_{\nu}|^2 \text{ endlich} \right)$$

definiert ist. $\bar{R}$ hat die Defektindizes 0, 1, es ist also max., aber nicht hypermax.

Beweis. Wurde soeben erbracht. —

Wir werden im folgenden noch einige wichtige allgemeine Eigenschaften von $\bar{R}$ herleiten, seiner expliziten Darstellung ist der Teil VI der „Bemerkungen“ gewidmet. Zunächst greifen wir auf unsere früheren Entwicklungen zurück und zeigen:

Satz 38. R sei ein max., aber nicht hypermax. H. O., also seine Defektindizes 0, n ($n = 1, 2, \dots, \infty$) oder $m, 0$ ($m = 1, 2, \dots, \infty$). Es sei

$p = n$ bzw. $= m$, wir können dann $\mathfrak{H}$ in $p + 1$ abg. lin. M. $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ zerlegen, die paarweise orth. sind und zusammen die abg. lin. M. $\mathfrak{H}$ aufspannen (für $p = \infty$ bricht die Folge $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ nicht ab), so daß

$\alpha)$ jedes $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ R reduziert, der in ihnen liegende Teil von R sei bzw. $R_1, \dots, R_p, S$;

$\beta)$ jedes $\mathfrak{M}_q$ ($q = 1, 2, \dots, p$) ist unendlichvieldimensional und ein Hilbertscher Raum, d. h. $\mathfrak{H}$ isomorph; und in ihnen sind alle R_q dem $\bar{R}$ isomorph bzw. alle R_q dem $-\bar{R}$ isomorph ($\bar{R}$ nach Satz 37).

$\gamma)$ Für $\mathfrak{N}$ bestehen drei Möglichkeiten: erstens kann es unendlichvieldimensional und ein Hilbertscher Raum, d. h. $\mathfrak{H}$ isomorph sein, dann ist S ein hypermax. H. O.; zweitens kann es N ($= 1, 2, \dots$)-dimensional, also ein N -dimensionaler (komplexer) Euklidischer Raum, sein, dann ist S ein gewöhnlicher endlichvieldimensionaler Hermitescher Operator; drittens kann es 0-dimensional, d. h. $= (0)$ sein, dann kann man es, zusammen mit S , fortlassen.

Diese $R_1, \dots, R_p, S$ bestimmen R nach Satz 21.

Beweis. Wurde für $m = 0$, $n > 0$ schon im ersten Teile dieses Kapitels erbracht (man beachte noch Satz 37), für $m > 0$, $n = 0$ gewinnt man ihn durch Ersetzen von R durch $-R$ (also Vertauschen von m, n). —

In endlichvieldimensionalen Räumen kann jeder Hermitesche Operator auf die Diagonalf orm gebracht werden, solange er > 1 Dimensionen hat ist er also reduzibel. In $\mathfrak{H}$ sind ebenfalls die hypermax. H. O. reduzibel, wovon man sich wie folgt überzeugt.

Sei R hypermax., $F(\lambda)$ seine Z. d. E. nach Satz 36. Nehmen wir zuerst an, es gäbe ein λ , so daß $F(\lambda) \neq 0, 1$ ist, also die abg. lin. M. $\mathfrak{M}$ mit $P_{\mathfrak{M}} = F(\lambda) \neq (0)$, $\mathfrak{H}$. Dieses $\mathfrak{M}$ reduziert nun R , denn es ist erstens

$$\begin{aligned} \int_{-\infty}^{\infty} \lambda'^2 d|E(\lambda') E(\lambda) f|^2 &= \int_{-\infty}^{\infty} \lambda'^2 d|E(\text{Min}(\lambda', \lambda)) f|^2 \quad (\text{vgl. Anm. } 53) \\ &= \int_{-\infty}^{\lambda} \lambda'^2 d|E(\lambda') f|^2, \end{aligned}$$

und dies ist endlich, wenn $\int_{-\infty}^{\infty} \lambda' d|E(\lambda') f|^2$ es ist, und zweitens gilt dann

$$\begin{aligned} (E(\lambda) R f, g) &= (R f, E(\lambda) g) = \int_{-\infty}^{\infty} \lambda' d(E(\lambda') f, E(\lambda) g) \\ &= \int_{-\infty}^{\infty} \lambda' d(E(\lambda) E(\lambda') f, g) = \int_{-\infty}^{\infty} \lambda' d(E(\lambda') E(\lambda) f, g) = (R E(\lambda) f, g) \end{aligned}$$

für alle g , d. h. es ist $E(\lambda) R f = R E(\lambda) f$.

Sind aber alle $E(\lambda) = 0$ oder 1, so gibt es (wegen α) in Def. 17) ein λ_0 , so daß für $\lambda < \lambda_0$ das erstere und für $\lambda > \lambda_0$ das letztere gilt, λ_0 ist endlich (wegen β), und für $\lambda = \lambda_0$ gilt das zweite (wegen γ). Daher ist

$$(Rf, g) = \int_{-\infty}^{\infty} \lambda' d(E(\lambda)f, g) = \lambda_0(f, g), \quad Rf = \lambda_0 f,$$

d. h. $\lambda_0 \cdot 1$ eine Forts. von R , also $R = \lambda_0 \cdot 1$. Dieses wird aber offenbar überhaupt durch jede abg. lin. M. reduziert.

Die Reduzibilität ist somit ein Rudiment der Diagonal-Transformierbarkeit und als solches mit der Möglichkeit der Eigenwertdarstellung aufs engste verbunden; Irreduzibilität einer Matrix bedeutet bereits das Auftreten der Ausnahmefälle der Elementarteilertheorie — was durch den Hermiteschen Charakter im Endlichvioldimensionalen, und wie wir sahen auch in § im Hypermax., ausgeschlossen wird. Unter diesem Gesichtspunkte ist nun die folgende Eigenschaft von $\bar{R}$ sehr beachtenswert:

Satz 39. *Ein max. H. O. ist dann und nur dann irreduzibel, wenn er (bei geeigneter Wahl des vollst. norm. orth. Systems $\psi_0, \psi_1, \psi_2, \dots$ aus Satz 37) mit $\bar{R}$ oder mit $-\bar{R}$ zusammenfällt (beides läßt sich nie beim selben H. O. erreichen).*

Beweis. Sei R max. und irreduzibel. Nach dem früher Gesagten ist es dann gewiß nicht hypermax., also tritt Satz 38 in Kraft. Da das dortige $\mathfrak{M}_1$ R reduziert, und $\mathfrak{M}_1 \neq 0$ ist, muß $\mathfrak{M}_1 = \mathfrak{S}$ sein — also $p = 1$ und $\mathfrak{N}$ fehlend. Dann ist $R = R_1$ also nach β) dortselbst eben $\bar{R}$ oder $-\bar{R}$ (bei richtiger Wahl von $\psi_0, \psi_1, \psi_2, \dots$). Da im einen Falle die Defektindizes von R 0, 1 sind, im anderen 1, 0, ist sicher beides unvereinbar.

Es bleibt noch übrig zu zeigen, daß $\bar{R}$ (und also auch $-\bar{R}$) irreduzibel ist — d. h. daß $\bar{U}$ irreduzibel ist (vgl. Satz 25). Möge also $\bar{U}$ durch $\mathfrak{M}$ reduziert werden, also auch durch $\mathfrak{N} = \mathfrak{S} - \mathfrak{M}$; jede dieser Mengen wird also durch das (überall sinnvolle) $\bar{U}$ auf eine Teilmenge von sich abgebildet: $\overline{\mathfrak{M}}, \overline{\mathfrak{N}}$. Wäre $\overline{\mathfrak{M}} = \mathfrak{M}, \overline{\mathfrak{N}} = \mathfrak{N}$, so enthielte der Wertevorrat von $\bar{U}$ $\mathfrak{M}$ und $\mathfrak{N}$, also auch $\mathfrak{M} + \mathfrak{N} = \mathfrak{S}$, was nicht der Fall ist. Also ist $\overline{\mathfrak{M}}$ oder $\overline{\mathfrak{N}}$ echtes Teil von $\mathfrak{M}$ oder $\mathfrak{N}$, $\mathfrak{M} - \overline{\mathfrak{M}}$ oder $\mathfrak{N} - \overline{\mathfrak{N}} \neq (0)$. Sei $f \neq 0$ ein Element von $\mathfrak{M} - \overline{\mathfrak{M}}$ bzw. $\mathfrak{N} - \overline{\mathfrak{N}}$. Dann ist es orth. zu $\overline{\mathfrak{M}}$ bzw. $\overline{\mathfrak{N}}$, und als Element von $\mathfrak{M}$ bzw. $\mathfrak{N}$ auch zu $\mathfrak{N}$ bzw. $\mathfrak{M}$ und $\mathfrak{N}$ bzw. $\overline{\mathfrak{M}}$: also zum ganzen Wertevorrat von $\bar{U}, \overline{\mathfrak{M}} + \overline{\mathfrak{N}}$. D. h. zu allen $\psi_1, \psi_2, \dots$, es ist also $= a\psi_0$ (mit $a \neq 0$). Folglich gehört auch ψ_0 zu $\mathfrak{M} - \overline{\mathfrak{M}}$ bzw. $\mathfrak{N} - \overline{\mathfrak{N}}$, also zu $\mathfrak{M}$ bzw. $\mathfrak{N}$. Aber $\mathfrak{M}, \mathfrak{N}$ enthalten auch ihre $\bar{U}$ -Bilder, das betreffende von ihnen enthält daher mit ψ_0 auch $\psi_1, \psi_2, \dots$ — also ganz $\mathfrak{S}$. Darum ist $\mathfrak{M} = \mathfrak{S}$ oder $\mathfrak{S} - \mathfrak{M} = \mathfrak{N} = \mathfrak{S}$, $\mathfrak{M} = (0)$, also alles bewiesen. —

Mit $\bar{R}$ ist auch $a\bar{R} + b1$ (a, b reell, $a \neq 0$) max. und irreduzibel, also (bis aufs Koordinatensystem) $= \bar{R}$ oder $= -\bar{R}$. Für $a > 0$ sind die Defektindizes 0, 1, also kommt dann $\bar{R}$ in Frage, für $a < 0$ sind sie 1, 0, also dann $-\bar{R}$.

Zum Schluß sei noch erwähnt, daß die Zerlegung des H. O. R im Satz 38 eine Art Umkehrung zuläßt, die die Konstruktion von abg. lin. H. O. R mit beliebig vorgegebenen Defektindizes m, n ermöglicht. Da $m = n = 0$ durch jeden hypermax. H. O. (also z. B. 1) realisiert wird, können wir $m + n > 0$ annehmen.

Dann gibt es $m + n$ unendlichvieldimensionale abg. lin. M. $\mathfrak{M}_1, \dots, \mathfrak{M}_m, \mathfrak{N}_1, \dots, \mathfrak{N}_n$, die zu je zweien orth. sind und zusammen $\mathfrak{H}$ aufspannen (als abg. lin. M.)⁵⁷⁾; jede ist dem $\mathfrak{H}$ isomorph, wir können darum in den m ersten $-\bar{R}$ und den n übrigen $\bar{R}$ realisieren: es seien die Operatoren $R_1, \dots, R_m, S_1, \dots, S_n$. Man überlegt sich leicht, daß diese, nach der Anweisung von Satz 21 zusammengefaßt, einen abg. lin. H. O. R ergeben, und daß R die Defektindizes m, n hat.

XI. Halbbeschränktheit.

Wir wollen einige Merkmale angeben, die jedem nicht hypermax. (aber max.!) H. O. abgehen müssen. (Die meisten Resultate dieses Satzes gewinnen wir übrigens bald auch auf anderem Wege.)

Satz 40. *Ein max., aber nicht hypermax. H. O. kann weder nach oben noch nach unten halbbeschränkt (also erst recht nicht beschränkt) sein, er kann weder überall sinnvoll sein noch alle Werte annehmen, und er kann in keinem vollst. norm. orth. System reell sein* (vgl. Anm.²⁰⁾).

Beweis. Für einen solchen Operator gibt es eine abg. lin. M. $\mathfrak{M}$, die ihn reduziert, so daß sein in $\mathfrak{M}$ gelegener Teil $\bar{R}$ oder $-\bar{R}$ ist (Satz 38), hätte er nun eine der vier zuerst genannten Eigenschaften, so müßte dieselbe auch $\bar{R}$ oder $-\bar{R}$ zukommen. Aber $\bar{R}$ (und darum auch $-\bar{R}$) besitzt keine einzige derselben, wie sich bei seiner genauen Diskussion in Anhang IV zeigen wird (im § 4). Mit der letzten steht es so: Wenn ein Operator in einem vollst. norm. orth. System im Sinne von Anm.²⁰⁾ reell ist, so besteht für ihn kein Unterschied zwischen i und $-i$ — also sind seine Defektindizes einander gleich. Bei einem max. und nicht hypermax. H. O. sind sie aber verschieden (vgl. Satz 33). —

⁵⁷⁾ Sei $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System, wir zerlegen es in $m + n$ Teilfolgen $\psi_1^1, \psi_2^1, \dots, \psi_1^m, \psi_2^m, \dots, \chi_1^1, \chi_2^1, \dots, \chi_1^n, \chi_2^n, \dots$ (es kann auch m oder $n = \infty$ sein), und setzen $\mathfrak{M}_1, \dots, \mathfrak{M}_m, \mathfrak{N}_1, \dots, \mathfrak{N}_n$ den von einer jeden dieser Teilfolgen aufgespannten abg. lin. M. gleich.

Wir gehen nun zur systematischen Untersuchung der halbbeschränkten Operatoren über.

Satz 41. *Ein H.O. R , der alle Werte annimmt, ist hypermax. Er nimmt jeden Wert genau einmal an.*

Beweis. Sei f ein Erweiterungselement von R , ihm zugeordnet sei f^* . Es gibt nach Annahme ein g mit sinnvollem $Rg = f^*$. Für alle h mit sinnvollem Rh gilt dann: $(f, Rh) = (f^*, h) = (Rg, h) = (g, Rh)$, und da Rh alle Werte annimmt, muß $f = g$ sein; also ist Rf sinnvoll. Daher ist R hypermax., wir haben aber mitbewiesen: aus $Rf = Rg$ folgt, indem wir $f^* = Rf$ setzen, $f = g$.

Satz 42. *Ein lin. H.O. R , für den stets $(f, Rf) \geq |f|^2$ gilt, kann zu einem definiten hypermax. H.O. fortgesetzt werden.*

Beweis. Zunächst bilden wir den abg. lin. H.O. $\tilde{R}$, auch für diesen gilt offenbar $(f, \tilde{R}f) \geq |f|^2$. Der Wertevorrat von $\tilde{R}$ sei $\mathfrak{M}$, $\mathfrak{M}$ ist offenbar eine lin. M. Es ist

$$|f|^2 \leq (f, \tilde{R}f) \leq |f| \cdot |\tilde{R}f|, \quad |\tilde{R}f| \geq |f|, \quad |\tilde{R}f - \tilde{R}g| = |\tilde{R}(f - g)| \geq |f - g|,$$

also folgt aus $\tilde{R}f = \tilde{R}g$ $f = g$, d. h. $\tilde{R}$ hat eine Inverse $\tilde{R}^{-1}$. $\tilde{R}^{-1}$ ist offenbar auch abg. lin., in $\mathfrak{M}$ definiert, und wegen $|f - g| \geq |\tilde{R}^{-1}f - \tilde{R}^{-1}g|$ stetig — also ist $\mathfrak{M}$ abg. Also können wir die abg. lin. M. $\mathfrak{N} = \mathfrak{S} - \mathfrak{M}$ bilden.

f gehöre zu $\mathfrak{N}$. Wenn $\tilde{R}g$ Sinn hat, so gehört es zu $\mathfrak{M}$, also ist $(f, \tilde{R}g) = 0$ — d. h. f ist Erweiterungselement, zu ihm gehört 0. Wenn also $\tilde{R}f$ Sinn hat, so ist $\tilde{R}f = 0$, $f = 0$. Der Definitionsbereich $\mathfrak{A}$ von R (eine lin. M.!) und $\mathfrak{N}$ sind also lin. unabh. Wir können daher einen neuen Operator S so definieren: Sf hat Sinn, wenn $f = g + h$ (g von $\mathfrak{A}$, h von $\mathfrak{N}$, die Zerlegung ist ja eindeutig!) ist, und zwar $= \tilde{R}g$. S ist Forts. von $\tilde{R}$, und ein H.O., denn wenn Sf, Sf' Sinn haben ($f = g + h$, $f' = g' + h'$, g, g' von $\mathfrak{A}$, h, h' von $\mathfrak{N}$), so ist:

$$(f, Sg) = (g + h, \tilde{R}g') = (g, \tilde{R}g') = (\tilde{R}g, g') = (\tilde{R}g, g' + h') = (Sf, g).$$

$\mathfrak{N}$ reduziert S . In $\mathfrak{N}$ ist ja nach Definition Sf immer sinnvoll, und zwar $= 0$; da Sf stets einem $\tilde{R}g$ gleich, also in $\mathfrak{M}$ gelegen ist, haben wir auch $P_{\mathfrak{N}} Sf = 0 = SP_{\mathfrak{N}} f$. Also reduziert auch $\mathfrak{M} = \mathfrak{S} - \mathfrak{N}$ S , die in $\mathfrak{M}, \mathfrak{N}$ liegenden Teile von S seien S' bzw. S'' . Der Wertevorrat von S ist $\mathfrak{M}$, wegen $Sf = P_{\mathfrak{M}} Sf = SP_{\mathfrak{M}} f$ nimmt S jeden seiner Werte schon in $\mathfrak{M}$ an — also hat S' auch den Wertevorrat $\mathfrak{M}$. Als ein-eindeutig lin. Bild des unendlichvieldimensionalen (weil überall dichten) $\mathfrak{A}$ ist es $\mathfrak{M}$ ebenfalls, also ist $\mathfrak{M}$ ein Hilbertscher Raum; darum ist Satz 41 anwendbar: S' ist hypermax., S'' ist sogar $= 0$.

Hieraus folgert man mühelos, daß S selbst hypermax. sein muß. Obendrein ist es definit, denn sei $f = g + h$, g von $\mathfrak{A}$, h von $\mathfrak{N}$:

$$(f, Sf) = (g + h, \tilde{R}g) = (g, \tilde{R}g) \geq |g|^2 \geq 0.$$

Damit ist alles bewiesen.

Satz 43. *Ein lin. H. O. R , für den stets $(f, Rf) \geq C \cdot |f|^2$ oder stets $\leq C \cdot |f|^2$ gilt (die Bedingungen der Halbbeschränktheit!) kann zu einem hypermax. H. O. S fortgesetzt werden, für den stets $(f, Sf) \geq C' \cdot |f|^2$ bzw. stets $\leq C' \cdot |f|^2$ gilt — dabei kann für C' jede Zahl $< C$ bzw. $> C$ gewählt werden⁵⁸).*

Beweis. Für $C' = 0$, $C = 1$ (und $>$ -Zeichen) ist dies Satz 42, und alle anderen Fälle können hierauf zurückgeführt werden, indem man $\frac{1}{C-C'} \cdot R - \frac{C'}{C-C'} \cdot 1$, $\frac{1}{C-C'} \cdot S - \frac{C'}{C-C'} \cdot 1$ statt R, S betrachtet.

XII. Pathologie der unbeschränkten Operatoren.

Zunächst einige Hilfssätze.

Satz 44. *Wenn R ein nach oben nicht halbbeschränkter lin. H. O. ist, so gibt es (für jedes $n = 1, 2, \dots$ und jedes C) eine n -dimensionale lin. M., in der Rf stets Sinn hat, und durchweg $(f, Rf) \geq C \cdot |f|^2$ ist.*

Beweis. Seien $f_1, \dots, f_k$ irgendwelche Elemente, für die $Rf_1, \dots, Rf_k$ sinnvoll ist. Wir zeigen zuerst: es muß ein φ geben, so daß φ zu allen $f_1, \dots, f_k$ orth. ist, $|\varphi| = 1$, $R\varphi$ sinnvoll, und $(\varphi, R\varphi) \geq D$. Nehmen wir nämlich das Gegenteil an. Da es nur auf die von den $f_1, \dots, f_k$ aufgespannte lin. M. ankommt, können wir sie „orthogonalisieren“ (vgl. Satz 8), d. h. durch die norm. orth. $\varphi_1, \dots, \varphi_l$ ($l \leq k$) ersetzen. Sei f beliebig, aber mit sinnvollem Rf . Dann ist $f = \sum_{\mu=1}^l a_\mu \varphi_\mu + b\varphi$, wo $|\varphi| = 1$, φ zu allen $\varphi_1, \dots, \varphi_l$ orth. ist (wir „orthogonalisieren“ f hinzu), da $Rf, R\varphi_1, \dots, R\varphi_l$ Sinn haben, hat auch $R\varphi$ Sinn — nach Annahme muß $(\varphi, R\varphi) \leq D$ sein. Es ist

$$|f|^2 = \sum_{\mu=1}^l |a_\mu|^2 + |b|^2,$$

$$\begin{aligned} (f, Rf) &= \sum_{\mu,\nu=1}^l a_\mu \bar{a}_\nu (\varphi_\mu, R\varphi_\nu) + 2\Re \sum_{\mu=1}^l b \bar{a}_\mu (\varphi, R\varphi_\mu) + |b|^2 (\varphi, R\varphi) \\ &\leq \sum_{\mu,\nu=1}^l |a_\mu| |a_\nu| |\varphi_\mu| |R\varphi_\nu| + 2 \sum_{\mu=1}^l |b| |a_\mu| |\varphi| |R\varphi_\mu| + |b|^2 \cdot D, \end{aligned}$$

⁵⁸) Der Satz gilt wohl auch noch für $C' = C$, es liegt aber kein Beweis dafür vor.

also wenn wir $D' = \text{Max}_{\mu=1, \dots, l} (|R \varphi_\mu|, D)$ setzen,

$$\leq D' \cdot \left(\sum_{\mu=1}^l |a_\mu| + |b| \right)^2 \leq (l+1) D' \cdot \left(\sum_{\mu=1}^l |a_\mu|^2 + |b|^2 \right) = (l+1) D' \cdot |f|^2;$$

also wäre R nach oben halbbeschränkt, entgegen der Annahme.

Wir gehen nun einen Schritt weiter. $g_1, \dots, g_k$ seien nunmehr ganz willkürlich, es gibt dann ein φ mit sinnvollem $R\varphi$, $|\varphi| = 1$, $(\varphi, R\varphi) \geq D$, so daß alle $|\varphi, g_\mu| \leq \varepsilon$ sind ($\mu = 1, \dots, k$, $\varepsilon > 0$). (Also nicht ganz genau orth.) Die f mit sinnvollem Rf liegen ja überall dicht, wir wählen k solche $f_1, \dots, f_k$ mit $|g_\mu - f_\mu| \leq \varepsilon$ ($\mu = 1, \dots, k$), und φ nach dem vorher Bewiesenen. Dann ist $|\varphi| = 1$, $(\varphi, R\varphi) \geq D$, und

$$|(\varphi, g_\mu)| = |(\varphi, g_\mu - f_\mu)| \leq |\varphi| \cdot |g_\mu - f_\mu| \leq \varepsilon.$$

Jetzt können wir die eigentliche Behauptung beweisen. Wir wählen nun sukzessiv: φ_1 mit sinnvollem $R\varphi_1$, $|\varphi_1| = 1$, $(\varphi_1, R\varphi_1) \geq D$; φ_2 mit sinnvollem $R\varphi_2$, $|\varphi_2| = 1$, $(\varphi_2, R\varphi_2) \geq D$ und $|(R\varphi_1, \varphi_2)| \leq 1$; ...; φ_n mit sinnvollem $R\varphi_n$, $|\varphi_n| = 1$, $(\varphi_n, R\varphi_n) \geq D$ und $|(R\varphi_1, \varphi_n)| \leq 1, \dots, |(R\varphi_{n-1}, \varphi_n)| \leq 1$. Also: alle $R\varphi_\mu$ haben Sinn, es ist $|\varphi_\mu| = 1$, $(\varphi_\mu, R\varphi_\mu) \geq D$, und $|(R\varphi_\mu, \varphi_\nu)| \leq 1$ für $\mu < \nu$, also immer für $\mu \neq \nu$. Daraus folgt:

$$\begin{aligned} \left| \sum_{\mu=1}^n a_\mu \varphi_\mu \right|^2 &= \sum_{\mu, \nu=1}^n a_\mu \bar{a}_\nu (\varphi_\mu, \varphi_\nu) \leq \sum_{\mu, \nu=1}^n |a_\mu| |a_\nu| |\varphi_\mu| |\varphi_\nu| \\ &\leq \left(\sum_{\mu=1}^n |a_\mu| \right)^2 \leq n \cdot \sum_{\mu=1}^n |a_\mu|^2, \\ \left(\sum_{\mu=1}^n a_\mu \varphi_\mu, R \left\{ \sum_{\mu=1}^n a_\mu \varphi_\mu \right\} \right) &= \sum_{\mu, \nu=1}^n a_\mu \bar{a}_\nu (\varphi_\mu, R\varphi_\nu) \\ &= \sum_{\mu=1}^n |a_\mu|^2 (\varphi_\mu, R\varphi_\mu) + \sum_{\substack{\mu, \nu=1 \\ (\mu \neq \nu)}}^n a_\mu \bar{a}_\nu (\varphi_\mu, R\varphi_\nu) \\ &\geq D \sum_{\mu=1}^n |a_\mu|^2 - \sum_{\substack{\mu, \nu=1 \\ (\mu \neq \nu)}}^n |a_\mu| |a_\nu| = (D+1) \sum_{\mu=1}^n |a_\mu|^2 - \left(\sum_{\mu=1}^n |a_\mu| \right)^2 \\ &\geq (D+1-n) \sum_{\mu=1}^n |a_\mu|^2. \end{aligned}$$

Für $D > n-1$ hat also die zweite Ungleichung die Konsequenz, daß

$\sum_{\mu=1}^n a_\mu \varphi_\mu = 0$ auch $\sum_{\mu=1}^n |a_\mu|^2 = 0$, $a_1 = \dots = a_n = 0$ nach sich zieht — d. h. die $\varphi_1, \dots, \varphi_n$ sind lin. unabh., die lin. M. der $\sum_{\mu=1}^n a_\mu \varphi_\mu$ ist n -dimensional. Weiter ist (nach beiden Ungleichungen)

$$\left(\sum_{\mu=1}^n a_\mu \varphi_\mu, R \left\{ \sum_{\mu=1}^n a_\mu \varphi_\mu \right\} \right) \geq \frac{D+1-n}{n} \left| \sum_{\mu=1}^n a_\mu \varphi_\mu \right|^2,$$

also gilt, wenn wir noch $D \geq nC + n - 1$ nehmen, in dieser lin. M. durchweg $(f, Rf) \geq C \cdot |f|^2$. Damit sind wir am Ziele.

Satz 45. *Wenn R ein nach oben nicht halbbeschränkter lin. H. O. ist, und $f_1, \dots, f_k$ irgendwelche Elemente, C eine beliebige Zahl, so gibt es ein φ mit sinnvollem $R\varphi$, das zu allen $f_1, \dots, f_k$ orth. ist, und für welches $|\varphi| = 1$, $(\varphi, R\varphi) \geq C$ gilt.*

Beweis⁵⁹⁾. Wir wählen die lin. M. des Satzes 44 $k + 1$ -dimensional, in ihr sind dann die (höchstens k unabhängige lin. Bedingungen darstellenden) Gleichungen $(f_1, \varphi) = 0, \dots, (f_k, \varphi) = 0$ durch ein $f \neq 0$ lösbar — durch Multiplizieren mit einer Zahl können wir $|f| = 1$ erreichen. Dieses f leistet alles Gewünschte.

Satz 46. *Ein nach oben nicht halbbeschränkter lin. H. O. R ist immer Forts. von einem definiten lin. H. O. S .*

Beweis. Es genügt eine überall dichte Folge $f_1, f_2, \dots$ anzugeben, so daß alle Rf_μ sinnvoll sind und stets

$$\left(\sum_{\mu=1}^n x_\mu f_\mu, R \left\{ \sum_{\mu=1}^n x_\mu f_\mu \right\} \right) \geq 0 \quad (n = 1, 2, \dots)$$

gilt. Dann ist nämlich der in der von $f_1, f_2, \dots$ aufgespannten lin. M. (überall dicht, also nicht abg.!) definierte und dort mit R übereinstimmende Operator S offenbar ein lin. H. O. und definit — und R ist Forts. von ihm.

Sei $g_1, g_2, \dots$ eine überall dichte Folge, so daß alle Rg_μ sinnvoll sind (d. h. $g_1, g_2, \dots$ sei eine im Definitionsbereiche von $R, \mathfrak{A}$, überall dichte Teilfolge von $\mathfrak{A}$, $\mathfrak{A}$ ist ja seinerseits überall dicht; so etwas gibt es, weil $\mathfrak{A}$ als Teil von $\mathfrak{H}$ separabel ist, vgl. Anm. ³³⁾), und $C_1, C_2, \dots$ eine Folge von Zahlen, über die wir noch näher verfügen werden. Wir wählen nun eine Folge $\varphi_1, \varphi_2, \dots$ auf Grund von Satz 45 mit den folgenden Eigenschaften aus: $R\varphi_1$ sinnvoll, $|\varphi_1| = 1$, $(\varphi_1, R\varphi_1) \geq C_1$; $R\varphi_2$ sinnvoll, $|\varphi_2| = 1$, $(\varphi_2, R\varphi_2) \geq C_2$ und $(R\varphi_1, \varphi_2) = 0$; $R\varphi_3$ sinnvoll, $|\varphi_3| = 1$, $(\varphi_3, R\varphi_3) \geq C_3$ und $(R\varphi_1, \varphi_3) = (R\varphi_2, \varphi_3) = 0$; Somit sind alle $R\varphi_\mu$ sinnvoll, $|\varphi_\mu| = 1$, $(\varphi_\mu, R\varphi_\mu) \geq C_\mu$ und $(R\varphi_\mu, \varphi_\nu) = 0$ für $\mu < \nu$, also immer für $\mu \neq \nu$. Wir setzen $f_1 = g_1 + \varphi_1$, $f_2 = g_2 + \frac{1}{2}\varphi_2$, $f_3 = g_3 + \frac{1}{3}\varphi_3, \dots$, und behaupten, daß die Folge $f_1, f_2, \dots$ (bei geeigneter Wahl der $C_1, C_2, \dots$) alles Gewünschte leistet.

Zunächst ist sie überall dicht: $|f_n - g_n| = \left| \frac{1}{n} \varphi_n \right| = \frac{1}{n} \rightarrow 0$, zu jedem f gibt es eine Teilfolge $g_{N_1}, g_{N_2}, \dots$ von $g_1, g_2, \dots$ mit $g_{N_n} \rightarrow f$, dann ist auch $f_{N_n} \rightarrow f$. Nun berechnen wir (f, Rf) für die Linearaggregate der $f_1, f_2, \dots$:

⁵⁹⁾ Bei sinnvollen $Rf_1, \dots, Rf_k$ hatten wir dies schon am Anfange des Beweises von Satz 44, aber wir müssen uns von dieser Annahme befreien.

$$\begin{aligned}
\left(\sum_{\mu=1}^n a_{\mu} f_{\mu}, R \left\{ \sum_{\mu=1}^n a_{\mu} f_{\mu} \right\} \right) &= \left(\sum_{\mu=1}^n a_{\mu} g_{\mu} + \sum_{\mu=1}^n \frac{a_{\mu}}{\mu} \varphi_{\mu}, \sum_{\mu=1}^n a_{\mu} R g_{\mu} + \sum_{\mu=1}^n \frac{a_{\mu}}{\mu} R \varphi_{\mu} \right) \\
&= \sum_{\mu, \nu=1}^n a_{\mu} \bar{a}_{\nu} (g_{\mu}, R g_{\nu}) + 2 \Re \sum_{\mu, \nu=1}^n \frac{a_{\mu} \bar{a}_{\nu}}{\mu} (\varphi_{\mu}, R g_{\nu}) + \sum_{\mu, \nu=1}^n \frac{|a_{\mu}|^2}{\mu^2} (\varphi_{\mu}, R \varphi_{\mu}) \\
&\quad + \sum_{\substack{\mu, \nu=1 \\ (\mu \neq \nu)}}^n \frac{a_{\mu} \bar{a}_{\nu}}{\mu \nu} (\varphi_{\mu}, R \varphi_{\nu}) \\
&\geq - \sum_{\mu, \nu=1}^n |a_{\mu}| |a_{\nu}| |g_{\mu}| |R g_{\nu}| - 2 \sum_{\mu, \nu=1}^n |a_{\mu}| |a_{\nu}| \cdot \frac{|\varphi_{\mu}| |R g_{\nu}|}{\nu} + \sum_{\mu=1}^n |a_{\mu}|^2 \cdot \frac{C_{\mu}}{\mu^2} \\
&= \sum_{\mu=1}^n \frac{C_{\mu}}{\mu^2} \cdot |a_{\mu}|^2 - \sum_{\mu, \nu=1}^n \left(|g_{\mu}| |R g_{\nu}| + \frac{|R g_{\nu}|}{\nu} \right) |a_{\mu}| |a_{\nu}|.
\end{aligned}$$

Wir setzen $|g_{\mu}| |R g_{\nu}| + \frac{|R g_{\nu}|}{\nu} = \bar{A}_{\mu \nu}$ (dies hängt von $g_1, g_2, \dots$ ab, aber nicht von den $f_1, f_2, \dots$ oder $\varphi_1, \varphi_2, \dots$, d. h. den $C_1, C_2, \dots$!) und schätzen den Subtrahenden ab:

$$\begin{aligned}
\sum_{\mu, \nu=1}^n \bar{A}_{\mu \nu} |a_{\mu}| |a_{\nu}| &= \sum_{\mu=1}^n \bar{A}_{\mu \mu} |a_{\mu}|^2 + \sum_{\substack{\mu, \nu=1 \\ (\mu < \nu)}}^n (\bar{A}_{\mu \nu} + \bar{A}_{\nu \mu}) |a_{\mu}| |a_{\nu}| \\
&\leq \sum_{\mu=1}^n \bar{A}_{\mu \mu} |a_{\mu}|^2 + \sum_{\substack{\mu, \nu=1 \\ (\mu < \nu)}}^n (2^{\nu-\mu-1} (\bar{A}_{\mu \nu} + \bar{A}_{\nu \mu}) |a_{\nu}|^2 + \frac{1}{2^{\nu-\mu}} |a_{\mu}|^2) \\
&= \sum_{\mu=1}^n \left\{ \bar{A}_{\mu \mu} + \sum_{\varrho=1}^{\mu-1} 2^{\mu-\varrho-1} (\bar{A}_{\mu \varrho} + \bar{A}_{\varrho \mu}) \right\} + \sum_{\varrho=\mu+1}^n \frac{1}{2^{\varrho-\mu}} \Big\} |a_{\mu}|^2 \\
&\leq \sum_{\mu=1}^n \left\{ \bar{A}_{\mu \mu} + 1 + \sum_{\sigma=0}^{\mu-2} 2^{\sigma} (\bar{A}_{\mu, \mu-\sigma-1} + \bar{A}_{\mu-\sigma-1, \mu}) \right\} |a_{\mu}|^2.
\end{aligned}$$

Indem wir den $\{ \}$ -Ausdruck kurz mit $\bar{B}_{\mu}$ bezeichnen, wird:

$$\left(\sum_{\mu=1}^n a_{\mu} f_{\mu}, R \left\{ \sum_{\mu=1}^n a_{\mu} f_{\mu} \right\} \right) \geq \sum_{\mu=1}^n \left(\frac{1}{\mu^2} C_{\mu} - \bar{B}_{\mu} \right) \cdot |a_{\mu}|^2.$$

Es genügt also $C_{\mu} = \mu^2 \cdot \bar{B}_{\mu}$ zu setzen und wir sind am Ziele⁶⁰⁾. —

⁶⁰⁾ Für die Zwecke der in Anm. ²³⁾ angekündigten Arbeit soll hier noch eine gewisse Verschärfung dieser Abschätzung ausgeführt werden (hier brauchen wir sie nicht). Es ist

$$\begin{aligned}
\left| \sum_{\mu=1}^n a_{\mu} f_{\mu} \right|^2 &= \left| \sum_{\mu=1}^n a_{\mu} g_{\mu} + \sum_{\mu=1}^n \frac{a_{\mu}}{\mu} \varphi_{\mu} \right|^2 \\
&= \sum_{\mu, \nu=1}^n a_{\mu} \bar{a}_{\nu} (g_{\mu}, g_{\nu}) + 2 \Re \sum_{\mu, \nu=1}^n \frac{a_{\mu} \bar{a}_{\nu}}{\mu} (\varphi_{\mu}, g_{\nu}) + \sum_{\mu, \nu=1}^n \frac{a_{\mu} \bar{a}_{\nu}}{\mu \nu} (\varphi_{\mu}, \varphi_{\nu}) \\
&\leq \sum_{\mu, \nu=1}^n \left\{ |g_{\mu}| |g_{\nu}| + \frac{|\varphi_{\mu}| |g_{\nu}|}{\mu} + \frac{|\varphi_{\mu}| |\varphi_{\nu}|}{\mu \nu} \right\} |a_{\mu}| |a_{\nu}| \\
&= \sum_{\mu, \nu=1}^n \left\{ |g_{\mu}| |g_{\nu}| + \frac{|g_{\nu}|}{\mu} + \frac{1}{\mu \nu} \right\} |a_{\mu}| |a_{\nu}|.
\end{aligned}$$

Damit haben wir bereits ein recht „pathologisches“ Resultat gewonnen, das sich leicht weiter verschärfen läßt.

Satz 47. *Sei R ein nach oben oder nach unten nicht halbbeschränkter lin. H. O., dann gibt es einen hypermax. H. O. S , für den stets $(f, Sf) \geq C \cdot |f|^2$ bzw. stets $\leq C \cdot |f|^2$ gilt, so daß R, S beide Forts. desselben lin. H. O. sind — dabei kann für C jede Zahl gewählt werden.*

Beweis. Für $C = -1$ und nicht-Halbbeschränktheit nach oben folgt dies aus Satz 46 und Satz 43 (wo $C = 0$, $C' = -1$ zu setzen ist), für alle anderen C folgt es hieraus, wenn man $R - (C+1) \cdot 1$, $S - (C+1) \cdot 1$ bzw. $-R + (C-1) \cdot 1$, $-S + (C-1) \cdot 1$ statt R, S betrachtet. —

Hieraus folgen die sonderbarsten Dinge für unbeschränkte H. O. Das Verhältnis eines H. O. zu seiner Matrix und dieser zu ihren Abschnitten wird dadurch qualitativ anders wie bei beschränkten; insbesondere zeigt sich, daß die „Abschnittsspektren“ bei unbeschränkten Matrizen sehr wenig mit dem Spektrum der Matrix selbst zu tun haben und diese (nicht wie in der Hilbertschen Theorie) in keinem Sinne „approximieren“. Über das Verhältnis Operator-Matrix werden wir einiges im Anhang III sagen, die genaue Erörterung der Eigenschaften unbeschränkter Matrizen soll in der in Anm. ²²) angekündigten Arbeit erfolgen.

Wir geben noch einige charakteristische Eigenschaften der unbeschränkten H. O. an. Im Gegensatze zu Satz 41 bis 47 beschränken wir uns jetzt auf abg. lin. Operatoren.

Satz 48. *Sei R ein abg. lin. H. O. Dann sind die vier Aussagen*
 $\alpha)$ R ist beschränkt, $\beta)$ R ist beschränkt hypermax., $\gamma)$ R ist überall

Wir nennen den $\{ \}$ -Ausdruck kurz $\bar{C}_{\mu\nu}$, dann wird wie im Text:

$$\sum_{\mu, \nu=1}^n \bar{C}_{\mu\nu} |a_\mu| |a_\nu| \leq \sum_{\mu=1}^n \left\{ \bar{C}_{\mu\mu} + 1 + \sum_{\sigma=0}^{\mu-2} 2^\sigma (\bar{C}_{\mu, \mu-\sigma-1} + \bar{C}_{\mu-\sigma-1, \mu}) \right\} |a_\mu|^2.$$

Diesen $\{ \}$ -Ausdruck nennen wir $\bar{D}_\mu$, dann wird

$$\left| \sum_{\mu=1}^n a_\mu f_\mu \right|^2 \leq \sum_{\mu=1}^n \bar{D}_\mu |a_\mu|^2.$$

Wenn wir nun $C_\mu = \mu^2 (\bar{B}_\mu + \mu \bar{D}_\mu)$ setzen, so ist für $a_1 = \dots = a_{p-1} = 0$ ($n \geq p$!)

$$\left(\sum_{\mu=p}^n a_\mu f_\mu, R \left\{ \sum_{\mu=p}^n a_\mu f_\mu \right\} \right) - p \left| \sum_{\mu=p}^n a_\mu f_\mu \right|^2 \geq \sum_{\mu=p}^n \left(\frac{C_\mu}{\mu^2} - \bar{B}_\mu - p \bar{D}_\mu \right) |a_\mu|^2 \geq 0.$$

D. h.: in der von den $f_p, f_{p+1}, f_{p+2}, \dots$ aufgespannten lin. M. (auch diese ist natürlich überall dicht, also nicht abg.!) gilt sogar

$$(f, Rf) \geq p \cdot |f|^2,$$

für jedes $p = 1, 2, \dots$

sinnvoll, δ) es gibt keinen abg. lin. H. O. von dem R eig. Forts. ist, alle gleichwertig⁶¹⁾).

Beweis. Wenn R nach oben nicht halbbeschränkt ist, so wählen wir C so, daß $(f, Rf) < C \cdot |f|^2$ vorkommt, und dann einen hypermax. H. O. S nach Satz 47, so daß stets $(f, Rf) \geq C \cdot |f|^2$ gilt — dann kann S offenbar keine Forts. von R sein. Nun sind R, S Forts. eines lin. H. O. T , und weil sie abg. lin. sind, auch des abg. lin. H. O. $\tilde{T}$. Da S nicht Forts. von R ist, muß $\tilde{T} \neq R$ sein, d. h. δ) ist nicht erfüllt. Wenn R nach unten nicht halbbeschränkt ist, so gilt δ) nicht für $-R$, also auch nicht für R . Also folgt allenfalls aus der Unbeschränktheit das Gegenteil von δ), d. h. aus δ) folgt α).

Aus α) folgt γ): denn R ist abg. und stetig, hat also einen abg. Definitionsbereich; da aber dieser überall dicht ist, muß er $= \mathfrak{H}$ sein. Aus γ) folgt δ): denn wäre R eig. Forts. des abg. lin. H. O. S , so wäre dieser nicht max., also zu mehreren max. H. O. festsetzbar (Zusatz zu Satz 35), z. B. zu $R' \neq R$; dann ist für alle f mit sinnvollen $R'f$ (wenn Sg Sinn hat)

$$(R'f, g) = (f, R'g) = (f, Sg) = (f, Rg) = (Rf, g),$$

und weil die Menge dieser g überall dicht ist, $R'f = Rf$ — also wäre R Forts. von R' , $R' \neq R$, trotzdem R' max. sein sollte.

Wir haben das logische Schema

$$\alpha) \rightarrow \gamma) \rightarrow \delta) \rightarrow \alpha),$$

d. h. α), γ), δ) sind logisch gleichwertig. Aus β) folgt α), und β) folgt aus α) und γ): daher ist es ihnen auch gleichwertig.

Satz 49. Sei R ein unbeschränkter abg. lin. H. O., dann ist es Forts. eines abg. lin. H. O., der zu Kontinuum vielen verschiedenen max. H. O. fortsetzbar ist.

Beweis. Nach Satz 48 δ) muß R eig. Forts. eines abg. lin. H. O. S sein, dieses S ist nicht max., hat also nach dem Zusatz zu Satz 35 die gewünschte Eigenschaft.

Anhang I. Funktionenräume.

Hier soll der im Kap. I angekündigte Nachweis erbracht werden, daß alle in der Einleitung I. und II. erwähnten Funktionenräume abstrakte Hilbertsche Räume sind, d. h. den Bedingungen A bis E aus Kap. I genügen. Für den „gewöhnlichen Hilbertschen Raum“ (aller Folgen $(x_1, x_2, \dots)$ mit $\sum_{n=1}^{\infty} |x_n|^2$) erübrigt sich dieser Beweis, sobald mindestens ein A bis E ge-

⁶¹⁾ Die Gleichwertigkeit von α) und γ) ist der bekannte Satz von Toeplitz, Math. Annalen 69 (1911), S. 321.

nügendes System bekannt ist — denn jedes solche System ist ja (nach den Schlußüberlegungen von Kap. I) dem „gewöhnlichen Hilbertschen Raume“ isomorph, also muß dieser die Eigenschaften A bis E gleichfalls besitzen. Wir müssen daher nur noch die Funktionenräume von Einleitung I. betrachten.

Ω sei also einer der dort genannten metrischen Räume, oder auch irgendein allgemeineres Gebilde, von dem wir aber jedenfalls folgendes verlangen: Es soll in Ω einen Begriff der „Meßbarkeit“, ein „Maß“ und ein „Integral“ ($\int_{\Omega} f(P) dv$ möge es heißen) geben, und zwar im Lebesgueschen Sinne. Wir wollen hier nicht näher ausführen, was dieser „Lebesguesche Sinn“ genau bedeuten soll, dies wäre leicht zu beschreiben, aber etwas weitläufig. Bei den nun auszuführenden Überlegungen wird man ohnehin erkennen, auf welche Eigenschaften dieser Begriffe es ankommt. Ferner setzen wir Ω auch als separabel voraus.

Der Funktionenraum $\mathfrak{S}'$ bestehe nunmehr aus allen in Ω definierten meßbaren (komplexwertigen) Funktionen $f(P)$, für die $\int_{\Omega} |f(P)|^2 dv$ endlich ist: dabei gelten zwei Funktionen, die nur auf einer Menge vom Maße 0 verschiedene Werte haben, als nicht verschieden. Was af und $f+g$ bedeuten, ist klar, (f, g) sei gleich $\int_{\Omega} f(P) \overline{g(P)} dv$. Diese Definitionen sind zulässig: wenn f, g zu $\mathfrak{S}'$ gehören, also $\int_{\Omega} |f(P)|^2 dv, \int_{\Omega} |g(P)|^2 dv$ endlich sind, so gehören auch $af, f+g$ zu $\mathfrak{S}'$, denn $\int_{\Omega} |af(P)|^2 dv$ ist offenbar endlich, und $\int_{\Omega} |f(P) + g(P)|^2 dv$ ist es wegen $|f(P) + g(P)|^2 \leq 2|f(P)|^2 + 2|g(P)|^2$; ferner ist $\int_{\Omega} f(P) \overline{g(P)} dv$ sinnvoll, ja absolut konvergent, wegen $|f(P) \overline{g(P)}| \leq \frac{1}{2}|f(P)|^2 + \frac{1}{2}|g(P)|^2$.

Wir gehen nun die Bedingungen A bis E der Reihe nach durch.

A, B sind offenbar erfüllt.

C. Es gilt eine Folge $f_1(P), f_2(P), \dots$ zu finden, die im Sinne der Entfernung

$$|f - g| = \sqrt{\int_{\Omega} |f(P) - g(P)|^2 dv}$$

überall dicht ist. — Das Maß von Ω ist eventuell unendlich, jedenfalls können wir aber Ω durch eine aufsteigende Folge $\Pi_1, \Pi_2, \dots$ von Teilmengen endlichen Maßes erschöpfen. Wenn f irgendeine Funktion aus $\mathfrak{S}'$ ist, so definieren wir f_N ($N=1, 2, \dots$) so:

$$f_N(P) = \begin{cases} f(P), & \text{wenn } P \text{ in } \Pi_N \text{ liegt und } |f(P)| \leq N \text{ ist,} \\ 0 & \text{sonst} \end{cases}$$

Da für jedes N $f_N(P) \rightarrow f(P)$ und alle $\int_{\Omega} |f_N(P)|^2 dv$ beschränkt ($\leq \int_{\Omega} |f(P)|^2 dv$) sind, so muß $\int_{\Omega} |f_N(P) - f(P)|^2 dv \rightarrow 0$, d. h. im Sinne unserer Entfernung $f_N \rightarrow f$. Also sind diejenigen f , die nur auf einer Menge endlichen Maßes Werte $\neq 0$ haben und die außerdem beschränkt sind, überall dicht — es genügt somit eine in ihnen überall dichte Folge zu finden.

Sei f eine Funktion dieser Klasse, das Maß der Menge wo sie $\neq 0$ ist, M , die obere Grenze ihres Absolutwertes A . Wir wählen im Intervalle $-A, A$ eine Einteilung $\varrho_1, \dots, \varrho_k$ rationaler Zahlen, die dieses mit einer Maschenweite $\leq \varepsilon$ überdecken ($\varepsilon > 0$, sonst beliebig). Es sei

$$f_{\varepsilon}(P) = \begin{cases} \varrho_h, & \text{wenn } f(P) \neq 0 \text{ und } \varrho_h \text{ das } f(P) \text{ nächstgelegene der } \varrho_1, \dots, \varrho_k \text{ ist,} \\ 0 & \text{sonst} \end{cases}$$

Hieraus folgt sofort $\int_{\Omega} |f_{\varepsilon}(P) - f(P)|^2 dv \leq M \cdot 2A \cdot \varepsilon$, für $\varepsilon \rightarrow 0$ gilt somit für diese f_{ε} $f_{\varepsilon} \rightarrow f$. D. h. diejenigen Funktionen, die stets rationale Werte annehmen, und nur endlich viele verschiedene, und jeden der $\neq 0$ ist nur auf einer Menge endlichen Maßes, sind auch überall dicht. Die gesuchte Folge braucht daher nur in diesen überall dicht zu sein.

Wenn Σ irgendeine Teilmenge von endlichem Maße von Ω ist, so sei f_{Σ} die Funktion, die in $\Sigma = 1$ und sonst $= 0$ ist; die Funktionenklasse von vorhin ist einfach die aller $\varrho_1 f_{\Sigma_1} + \dots + \varrho_k f_{\Sigma_k}$ ($\varrho_1, \dots, \varrho_k$ rational). Wenn wir eine Folge von Mengen $\Sigma^{(1)}, \Sigma^{(2)}, \dots$ fänden, die in der Gesamtheit aller Mengen endlichen Maßes überall dicht ist (die Entfernung von zwei Mengen Σ', Σ'' ist die ihrer Funktionen $f_{\Sigma'}, f_{\Sigma''}$, d. h. das Maß ihrer „Unterschiedsmenge“ $(\Sigma' + \Sigma'') - (\Sigma' \cdot \Sigma'')$ ⁶²⁾), so lägen die Funktionen $\varrho_1 f_{\Sigma^{(n_1)}} + \dots + \varrho_k f_{\Sigma^{(n_k)}}$ ($k = 1, 2, \dots, \varrho_1, \dots, \varrho_k$ rational, $(n_1, \dots, n_k = 1, 2, \dots)$ in $\mathfrak{S}'$ überall dicht: d. h. wir wären am Ziele. Eine solche Mengenfolge $\Sigma^{(1)}, \Sigma^{(2)}, \dots$ gilt es also zu finden.

Sei Σ eine Menge vom endlichen Maß M , bekanntlich ist M die untere Grenze der Maße aller Σ enthaltenden offenen Punktmengen; sei Σ' eine solche Menge mit einem Maße $\leq M + \varepsilon$, dann hat die „Unterschiedsmenge“ $\Sigma' - \Sigma$ ein Maß $\leq \varepsilon$. D. h. die offenen Mengen liegen überall dicht. Da Ω separabel ist, gibt es in Ω eine Folge offener Mengen $A_1, A_2, \dots$, die ein „gleichwertiges Umgebungssystem“ bilden, also so, daß jede offene Menge Σ' die Vereinigungsmenge aller in ihr enthaltenen A_n ist — etwa von $A_{n_1}, A_{n_2}, \dots$. Das Maß von Σ' sei M' , dann konvergieren die Maße der $A_{n_1}, A_{n_1} + A_{n_2}, A_{n_1} + A_{n_2} + A_{n_3}, \dots$ gegen M' . Sei etwa das-

⁶²⁾ $\pm, \cdot$ sind hier die Zeichen für das Vereinigen, Abziehen und Durchschnittsbilden bei Mengen.

jenige von $\Lambda_{n_1} + \dots + \Lambda_{n_k}$ bereits $\geq M' - \varepsilon$, dann hat die „Unterschiedsmenge“ $\Sigma' - (\Lambda_{n_1} + \dots + \Lambda_{n_k})$ ein Maß $\leq \varepsilon$. D. h.: die $\Lambda_{n_1} + \dots + \Lambda_{n_k}$ ($k = 1, 2, \dots, n_1, \dots, n_k = 1, 2, \dots$) liegen auch überall dicht⁶³), und diese sind abzählbar viel, womit wir am Ziele sind.

D. Sei $k = 1, 2, \dots$, wir wählen k paarweise punktfremde Mengen $K_1, \dots, K_k$ in Ω aus, jede von endlichem Maße $\neq 0$. Die Funktion $f_h(P)$ sei $= 1$ in K_h und $= 0$ sonst ($h = 1, \dots, k$); dann sind die $f_1, \dots, f_k$ offenbar. lin. unabh. (und sie gehören zu $\mathfrak{H}'$).

E. Sei $f_1, f_2, \dots$ eine Funktionenfolge, die dem Cauchyschen Konvergenzkriterium genügt (zu jedem $\varepsilon > 0$ gibt es ein $N = N(\varepsilon)$, so daß für $m, n \geq N$ $\int_{\Omega} |f_m(P) - f_n(P)|^2 dv \leq \varepsilon$ ist). Wir wählen (in Anlehnung an einen Beweis von Weyl) eine Folge $N_1 < N_2 < N_3 < \dots$ mit $N_p \geq N\left(\frac{1}{8^p}\right)$ aus und schließen dann so. Es ist

$$\int_{\Omega} |f_{N_{p+1}}(P) - f_{N_p}(P)|^2 dv \leq \frac{1}{8^p},$$

also hat diejenige Menge, auf der

$$|f_{N_{p+1}}(P) - f_{N_p}(P)| \geq \frac{1}{2^p}$$

gilt, ein Maß $\leq \frac{1}{2^p}$. Also gilt

$$|f_{N_{p+1}}(P) - f_{N_p}(P)| \leq \frac{1}{2^p}, \quad |f_{N_{p+2}}(P) - f_{N_{p+1}}(P)| \leq \frac{1}{2^{p+1}}, \dots$$

überall, mit Ausnahme einer Menge vom Maße $\leq \frac{1}{2^p} + \frac{1}{2^{p+1}} + \dots = \frac{1}{2^{p-1}}$. Aber dort, wo die obigen Ungleichheiten gelten, hat die Folge $f_{N_1}(P), f_{N_2}(P), \dots$ einen Limes; also überall mit Ausnahme einer Menge vom Maße $\leq \frac{1}{2^{p-1}}$. Da dies für alle p gilt, hat die Ausnahmемenge das Maß 0.

Wir definieren nun: $f(P) = \lim f_{N_p}(P)$, soweit dieser Limes existiert, d. h. überall bis auf eine 0-Menge, und sonst $= 0$. Wenn nun $m \geq N(\varepsilon)$ ist, so gilt für fast alle p (nämlich sobald $N_p \geq N(\varepsilon)$ ist) $\int_{\Omega} |f_m(P) - f_{N_p}(P)|^2 dv \leq \varepsilon$. Also dürfen wir $p \rightarrow \infty$ ausführen, und es wird: $\int_{\Omega} |f_m(P) - f(P)|^2 dv \leq \varepsilon$. Somit gehört erstens $f_m - f$, also auch f , zu $\mathfrak{H}'$, und zweitens ist $f_m \rightarrow f$. D. h. alles ist bewiesen.

Anhang II. Das Eigenwertproblem unitärer Operatoren.

1. Wir betrachten im folgenden durchweg in ganz $\mathfrak{H}$ sinnvolle und beschränkte (stetige) lin. Operatoren R , auf die also Satz 12 α), β) an-

⁶³) Wir nehmen natürlich nur die von endlichem Maße.

wendbar ist. Der Hermitesche Charakter hingegen wird nicht immer Voraussetzung sein. Bei festem f hat $L(g) = (f, Rg)$ die Eigenschaften

$$L(a_1 g_1 + \dots + a_n g_n) = \bar{a}_1 L(g_1) + \dots + \bar{a}_n L(g_n),$$

$$|L(g)| \leq C \cdot |f| |g| = D \cdot |g|,$$

so daß der Satz von F. Riesz (vgl. Anm. ⁵²) anwendbar ist: $L(g) = (f^*, g)$. Wir nennen das so bestimmte $f^* = R^* f$, R^* ist ein überall sinnvoller und offenbar lin. Operator. Es hat die Eigenschaften

$$(f, Rg) = (R^* f, g), \quad (Rf, g) = (f, R^* g)$$

(die erste war Definition, die zweite folgt durch Vertauschen von f, g und nehmen der komplex-Konjugierten). Insbesondere ist

$$|R^* f|^2 = (R^* f, R^* f) = (f, R R^* f) \leq C \cdot |f| |R^* f|, \quad |R^* f| \leq C \cdot |f|,$$

d. h. auch R^* beschränkt⁶⁴).

Man verifiziert sofort $R^{**} = R$, $(R \pm S)^* = R^* \pm S^*$, $(RS)^* = S^* R^*$,⁶⁵ $(aR)^* = \bar{a} R^*$, $0^* = 0$, $1^* = 1$. Für H. O. ist $R = R^*$ charakteristisch (wenn sie beschränkt sind). Jeder unitäre Operator U gehört zu unserer Klasse (er ist wegen $|Uf| = |f|$ beschränkt), und ist in ihr durch $UU^* = U^*U = 1$, d. h. $U^{-1} = U^*$ gekennzeichnet⁶⁶).

2. Jetzt sei R sogar H. O., wir nehmen zuerst sein Eigenwertproblem nach der Methode von F. Riesz in Angriff. Sei $p(x) = \sum_{v=1}^n a_v x^v$ ein Polynom (die a_v reell), dann ist $p(R) = \sum_{v=0}^n a_v R^v$ ($R^0 = 1$) ein H. O. unserer Klasse; das Addieren, Subtrahieren, Multiplizieren dieser $p(R)$ ist offenbar ebenso zu vollziehen, wie das der $p(x)$ (also sind alle vertauschbar). Wir wählen C mit $|Rf| \leq C \cdot |f|$ (für alle f).

Satz 1*. Wenn für $-C \leq x \leq C$ $p(x) \geq 0$ ist, so ist $p(R)$ definit.

Beweis. Sei $\mathfrak{M}$ irgendeine endlichvioldimensionale (also abg.) lin. M., etwa durch das norm. orth. System $\varphi_1, \dots, \varphi_k$ aufgespannt. $R' = P_{\mathfrak{M}} R P_{\mathfrak{M}}$ ist ein H. O. mit $|R'f| = |P_{\mathfrak{M}} R P_{\mathfrak{M}} f| \leq |R P_{\mathfrak{M}} f| \leq C \cdot |P_{\mathfrak{M}} f| \leq C \cdot |f|$, d. h. von unserer Klasse, der durch $\mathfrak{M}$ offenbar reduziert wird (sein in

⁶⁴) Bei Matrizen ist dieses $*$ das Nehmen der konjugiert-transponierten, ebenso bei Integraloperator-Kernen, bei Differentialoperatoren die adjungierte.

⁶⁵) Es handelt sich um überall sinnvolle Operatoren, daher sind die $+$, $-$, $\cdot$ ohne weiteres sinnvoll (vgl. Anm. ³⁷).

⁶⁶) Wenn U unitär ist, so hat es eine Inverse U^{-1} , und es ist $(f, U^{-1}g) = (Uf, U U^{-1}g) = (Uf, g) = (f, U^*g)$ für alle f, g , also $U^{-1} = U^*$. Wenn $U^{-1} = U^*$ ist, so ist $(Uf, Ug) = (f, U^*Ug) = (f, g)$, also für $f = g$ $|Uf| = |f|$, so daß U längentreu ist; da es überall Sinn hat und alle Werte annimmt (U^{-1} hat ja überall Sinn!), ist es unitär.

$\mathfrak{H} - \mathfrak{M}$ liegender Teil ist 0!). Sein in $\mathfrak{M}$ liegender Teil ist also ein k -dimensionaler H. O., etwa mit den Eigenwerten $\lambda_1, \dots, \lambda_k$; wegen $|R'f| \leq C \cdot |f|$ sind alle $|\lambda_1|, \dots, |\lambda_k| \leq C$. Auch $p(R')$ wird durch $\mathfrak{M}$ reduziert, und sein in $\mathfrak{M}$ liegender Teil hat die Eigenwerte $p(\lambda_1), \dots, p(\lambda_k)$ — also lauter Zahlen ≥ 0 , d. h. er ist definit. Also: Wenn f in $\mathfrak{M}$ liegt so ist

$$(f, p(R')f) \geq 0.$$

Sei nun f beliebig, wir wählen $\mathfrak{M}$ als die von $f, Rf, \dots, R^n f$ aufgespannte (abg.) lin. M., dann ist f in $\mathfrak{M}$, $R'f = Rf, \dots, R'^n f = R^n f$, also $p(R')f = p(R)f$. Also $(f, p(R)f) \geq 0$; und da f ganz beliebig war, ist $p(R)$ definit.

Satz 2*. Wenn für $-C \leq x \leq C$ $|p(x)| \leq \varepsilon$ ist, so ist $|p(R)f| \leq \varepsilon \cdot |f|$.

Beweis. Satz 1* ist auf $\varepsilon \pm p(x)$ anwendbar, also $(f, \{\varepsilon \cdot 1 \pm p(R)\}f) \geq 0$, $|(f, p(R)f)| \leq \varepsilon \cdot |f|^2$. Nach Satz 12 (Gleichwertigkeit von γ) und α !) bedeutet das $|p(R)f| \leq \varepsilon \cdot |f|$.

Satz 3*. Wenn $P(x)$ ein im Intervalle $-C \leq x \leq C$ stetige Funktion ist, und $p_1(x), p_2(x), \dots$ eine dort gleichmäßig gegen $P(x)$ strebende Polynomfolge, so hat $p_n(R)f$ einen Limes f^* , der nur von R, f und $P(x)$ abhängt.

Beweis. Klar nach Satz 2*. —

Durch $P(R)f = f^*$ definieren wir also einen überall sinnvollen Operator, der lin. H. O. ist, weil es die $p_n(R)$ sind, und beschränkt nach Satz 48 (α) und (γ), oder auch weil es die $p_n(x)$ gleichmäßig sind (also auch die $p_n(R)$, Satz 2*).

Satz 4*. Die $P(R)$ sind ebenso zu addieren, subtrahieren, multiplizieren wie die $P(x)$, Satz 1*, 2* gilt für sie auch; wenn $P(x)$ ein Polynom ist, stimmt die alte und neue Definition überein.

Beweis. Die letzte Behauptung folgt daraus, daß $p_1(x) = p_2(x) = \dots = P(x)$ gesetzt werden kann; die übrigen gelten für die $p(x)$, also aus Stetigkeitsgründen auch für $P(x)$.

3. Nun biegen wir vom F. Rieszschen Wege ab.

Satz 5*. Sei $\mathfrak{M}_R$ die Menge aller f mit $Rf = 0$, $\mathfrak{N}_R$ die Menge aller Rf und ihrer Häufungspunkte. Dann ist $\mathfrak{M}_R^\perp = \mathfrak{H} - \mathfrak{N}_R^\perp$.

Beweis. Die Rf bilden eine lin. M., es genügt zu zeigen, daß $\mathfrak{M}_R$ aus allen zu dieser lin. M. orth. g besteht. Nun bedeutet das orth. stehen auf diese, $(g, Rf) = 0$ für alle f , d. h. $(Rg, f) = 0$ — also $Rg = 0$. Damit sind wir fertig.

Satz 6*. Wenn R, S definit sind, so ist $\mathfrak{M}_{R+S}$ Teil von $\mathfrak{M}_R$, und $\mathfrak{N}_R$ Teil von $\mathfrak{N}_{R+S}$.

Beweis. Aus $(R + S)g = 0$ folgt $0 \leq (g, Rg) \leq (g, (R + S)g) = 0$, $(g, Rg) = 0$. Nun gilt für definite R

$$|(f, Rg)|^2 \leq (f, Rf) \cdot (g, Rg)$$

(man beweist dies ebenso, wie die entsprechende Gleichung für $R = 1$ bei Satz 1), also hat $(g, Rg) = 0$ die Folge $(f, Rg) = 0$ (für alle f !), $Rg = 0$. Die obige Relation bedeutet daher: $\mathfrak{M}_{R+S}$ Teil von $\mathfrak{M}_R$. Nach Satz 5* ist also $\mathfrak{N}_{\tilde{R}}$ Teil von $\mathfrak{N}_{R+S}$.

Wir definieren:

$$f_a(x) = \text{Max}(x - a, 0), \quad \mathfrak{M}(R, a) = \mathfrak{M}_{f_a(R)}, \quad \mathfrak{N}(R, a) = \mathfrak{N}_{f_a(R)}.$$

Satz 7*. In $\mathfrak{M}(R, a)$ bzw. $\mathfrak{N}(R, a)$ gilt durchweg

$$(f, Rf) \leq a \cdot |f|^2 \quad \text{bzw.} \quad \geq a \cdot |f|^2.$$

Beweis. Sei $g_a(x) = \text{Max}(-(x - a), 0)$, da $f_a(x)g_a(x) = 0$, $f_n(x) + g_n(x) = x - a$ ist, haben wir für $S' = f_a(R)$, $S'' = g_a(R)$ nach Satz 4*. $S'S'' = S''S' = 0$, $S' - S'' = R - a \cdot 1$. Da stets $f_a(x) \geq 0$, $g_a(x) \geq 0$ ist, sind S', S'' auch definit. In $\mathfrak{M}(R, a) = \mathfrak{M}_{S'}$ ist $S'f = 0$, also $Rf = af - S''f$, $(f, Rf) = a \cdot |f|^2 - (f, S''f) \leq a \cdot |f|^2$. In $\mathfrak{N}(R, a) = \mathfrak{N}_{S'}$ befinden wir uns (nach Satz 5* und $S''S' = 0$) in $\mathfrak{M}_{S''}$, also ist $S''f = 0$, und $Rf = af + S'f$, $(f, Rf) = a \cdot |f|^2 + (f, S'f) \geq a \cdot |f|^2$.

Satz 8*. Wenn $a \leq b$ ist, so ist $\mathfrak{M}(R, a)$ Teil von $\mathfrak{M}(R, b)$ und $\mathfrak{N}(R, b)$ Teil von $\mathfrak{N}(R, a)$, ferner ist für $a < -C$ bzw. $> C$ $\mathfrak{M}(R, a) = (0)$ bzw. $= \mathfrak{S}$ und $\mathfrak{N}(R, a) = \mathfrak{S}$ bzw. $= (0)$.

Beweis. Die erste Behauptung folgt, da $f_a(x) \geq 0$, $f_b(x) - f_a(x) \geq 0$ ist, also $f_a(R), f_b(R) - f_a(R)$ definit, aus Satz 6*. Die zweite beweisen wir so: Gehörte ein $f \neq 0$ für $a < -C$ bzw. $> C$ zu $\mathfrak{M}(R, a)$ bzw. $\mathfrak{N}(R, a)$, so wäre

$$(f, Rf) \leq a \cdot |f|^2 < -C \cdot |f|^2 \quad \text{bzw.} \quad \geq a \cdot |f|^2 > C \cdot |f|^2,$$

entgegen der Definition von C (Satz 12 (γ)), also ist $\mathfrak{M}(R, a)$ bzw. $\mathfrak{N}(R, a) = (0)$. Also $\mathfrak{N}(R, a)$ bzw. $\mathfrak{M}(R, a) = \mathfrak{S} - (0) = \mathfrak{S}$.

Satz 9*. Es gilt ganz allgemein

$$(Rf, g) = \int a \, d(P_{\mathfrak{M}(R, a)} f, g)$$

(wobei das Integral über irgend zwei Grenzen $< -C, > C$ zu erstrecken ist).

Beweis. Für $f = g$ folgt es aus Satz 7*, 8*⁶⁷⁾. Durch Einsetzen von

⁶⁷⁾ R ist mit $f_a(R)$ vertauschbar (Satz 4*), wird also durch $\mathfrak{M}(R, a)$ reduziert (aus $f_a(R)f = 0$ folgt $f_a(R) \cdot Rf = R \cdot f_a(R)f = 0$, und Satz 20), d. h. es ist mit $P_{\mathfrak{M}(R, a)}$ vertauschbar.

Nun sei $x_0 < x_1 < \dots < x_p$, $x_0 < -C$, $x_p > C$, dann ist

$\frac{f+g}{2}, \frac{f-g}{2}$ und Subtrahieren entsteht die Gleichung allgemein für die Realteile; durch Ersetzen von f, g durch if, g auch für die Imaginärteile. —

Satz 9* ist im wesentlichen das Hilbertsche Resultat für beschränkte H. O. Wir gehen weiter und betrachten zwei H. O. R, S unserer Klasse.

Satz 10*. Wenn R, S vertauschbar sind, so sind es auch $P_{\mathfrak{M}(R, a)}, P_{\mathfrak{M}(S, b)}$.

Beweis. R, S sind vertauschbar, also auch alle $P(R), Q(S)$ (Polynome offenbar, andere Funktionen aus Stetigkeitsgründen) — also insbesondere $f_a(R), f_b(S)$. Daher reduziert $\mathfrak{M}(R, a) f_b(S)$ (vgl. die Überlegung in Anm. ⁶⁷), und $\mathfrak{N}(R, a) = 1 - \mathfrak{M}(R, a)$ ebenfalls. Darum muß $\mathfrak{M}(S, b)$ so gebildet werden, daß man die f mit $f_b(S)f = 0$ zuerst in $\mathfrak{M}(R, a)$ und dann in $\mathfrak{N}(R, a)$ bildet, und dann beide abg. lin. M. addiert. Hieraus schließt man mühelos die Vertauschbarkeit von $P_{\mathfrak{M}(R, a)}$ und $P_{\mathfrak{M}(S, b)}$.

4. Jetzt wenden wir uns den sog. normalen Operatoren A unserer Klasse zu, die durch $AA^* = A^*A$ charakterisiert sind. Die H. O. $R = \frac{1}{2}(A + A^*), S = \frac{1}{2i}(A - A^*)$ sind dann vertauschbar, und es ist $A = R + iS, A^* = R - iS$. Wir bilden nun die P. O.

$$E(A, a, b) = P_{\mathfrak{M}(R, a)} \cdot P_{\mathfrak{M}(S, b)}.$$

Es gilt der Satz:

Satz 11*. ⁶⁸) A und die $E(A, a, b)$ erfüllen die folgenden Bedingungen:

$$\begin{aligned} (Rf, f) &= \sum_{\nu=1}^p (R(P_{\mathfrak{M}(R, x_\nu)} - P_{\mathfrak{M}(R, x_{\nu-1})})f, f) = \sum_{\nu=1}^p (R P_{\mathfrak{M}(R, x_\nu) - \mathfrak{M}(R, x_{\nu-1})} f, f) \\ &= \sum_{\nu=1}^p (R P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} f, f) = \sum_{\nu=1}^p (R P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} f, P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} f). \end{aligned}$$

(Das letztere, weil R mit $P_{\mathfrak{M}(R, x_\nu)} - P_{\mathfrak{M}(R, x_{\nu-1})} = P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})}$ vertauschbar ist, also $P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} R P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} = R P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})}$ gilt.) Nach Satz 7* ist also der einzelne Addend in $\sum_{\nu=1}^p$

$$\begin{cases} \leq x_\nu & |P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} f|^2, \\ \geq x_{\nu-1} & |P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} f|^2, \end{cases}$$

$$|P_{\mathfrak{M}(R, x_\nu) \cdot \mathfrak{N}(R, x_{\nu-1})} f|^2 = |P_{\mathfrak{M}(R, x_\nu) - \mathfrak{M}(R, x_{\nu-1})} f|^2 = |P_{\mathfrak{M}(R, x_\nu)} f|^2 - |P_{\mathfrak{M}(R, x_{\nu-1})} f|^2.$$

Nach Definition des Stieltjesschen Integrals bedeutet aber

$$(Rf, f) \begin{cases} \leq \sum_{\nu=1}^p x_\nu (|P_{\mathfrak{M}(R, x_\nu)} f|^2 - |P_{\mathfrak{M}(R, x_{\nu-1})} f|^2), \\ \geq \sum_{\nu=1}^p x_{\nu-1} (|P_{\mathfrak{M}(R, x_\nu)} f|^2 - |P_{\mathfrak{M}(R, x_{\nu-1})} f|^2), \end{cases}$$

das aus unseren Relationen folgt, die behauptete Gleichung

$$(Rf, f) = \int a d|P_{\mathfrak{M}(R, a)} f|^2 = \int a d(P_{\mathfrak{M}(R, a)} f, f).$$

⁶⁸) Vgl. hierzu Einleitung IX und Anm. ²⁶).

$\alpha)$ $E(A, a, b)$, $E(A, a', b')$ sind vertauschbar und ihr Produkt ist $E(A, \text{Min}(a, a'), \text{Min}(b, b'))$.

$\beta)$ Für genügend kleines a oder b ist $E(A, a, b) = 0$. Für genügend großes a oder b ist es von a bzw. b unabhängig. Für genügend großes a und b ist es $= 1$.

Und es gelten die Gleichungen:

$$(Af, g) = \iint (a + ib) d(E(A, a, b) f, g),$$

$$(A^* f, g) = \iint (a - ib) d(E(A, a, b) f, g).$$

(Die $\iint$ sind über ein hinreichend großes Rechteck zu erstrecken.)

Beweis. Vergegenwärtigt man sich die aus Satz 9*, 10* folgenden Gleichungen

$$(Rf, g) = \iint a d(E(A, a, b) f, g), \quad (Sf, g) = \iint b d(E(A, a, b) f, g),$$

so folgt alles unmittelbar aus Satz 8*, 9*, 10* und der Definition der $E(A, a, b)$. —

Die unitären Operatoren U sind wegen $UU^* = U^*U = 1$ auch normal, Satz 11* gilt also für sie auch. Ihre $E(U, a, b)$ haben aber noch weitergehende Eigenschaften und diese wollen wir jetzt herleiten.

Zunächst gilt allgemein:

$$\begin{aligned} (Af, E(A, a, b) g) &= \iint (a' + ib') d(E(A, a', b') f, E(A, a, b) g) \\ &= \iint (a' + ib') d(E(A, a, b) E(A, a', b') f, g) \\ &= \iint (a' + ib') d(E(A, \text{Min}(a, a'), \text{Min}(b, b')) f, g) \\ &= \int_a^a \int_b^b (a' + ib') d(E(A, a', b') f, g). \end{aligned}$$

$$\begin{aligned} (Af, Ag) &= \iint (a - ib) d(Af, E(A, a, b) g) \\ &= \iint (a - ib) d\left\{ \int_a^a \int_b^b (a' + ib') d(E(A, a', b') f, g) \right\} = (\text{vgl. Anm. }^{55}) \\ &= \iint (a - ib)(a + ib) d(E(A, a, b) f, g) \\ &= \iint (a^2 + b^2) d(E(A, a, b) f, g). \end{aligned}$$

$$(f, g) = \iint d(E(A, a, b) f, g).$$

Für unitäre A ist also wegen $(f, g) = (Af, Ag)$

$$\iint (1 - a^2 - b^2) d(E(A, a, b) f, g) = 0.$$

Und wenn wir $f = g$ setzen: $\iint (1 - a^2 - b^2) d|E(A, a, b) f|^2 = 0$.

5. Wenn $\Re$ irgendein Rechteck $a' \leq x \leq a''$, $b' \leq y \leq b''$ ist (Ecken: $a', b'; a'', b'; a', b''; a'', b''$), so definieren wir einen P. O.

$$\begin{aligned}
E(A, \Re) &= E(A, a'', b'') - E(A, a'', b') - E(A, a', b'') + E(A, a', b') \\
&= P_{\Re(R, a')} \cdot (1 - P_{\Re(R, a'')}) \cdot P_{\Re(S, b')} \cdot (1 - P_{\Re(S, b'')}) \\
&= P_{\Re(R, a') \Re(R, a'') \Re(S, b') \Re(S, b'')}.
\end{aligned}$$

Möge nun $\Re$ keinen einzigen Punkt mit dem Einheitskreis gemeinsam haben, dann ist in $\Re$ stets $1 - a^2 - b^2 \geq \varepsilon$ oder stets $\leq -\varepsilon$ ($\varepsilon > 0$), und wir haben für jedes f mit $E(A, \Re)f = f$ (d.h. aus $\Re(R, a') \Re(R, a'') \Re(S, b') \Re(S, b'')$)

$$\begin{aligned}
& \left| \iint (1 - a^2 - b^2) d |E(A, a, b) f|^2 \right| \\
&= \left| \iint_{\Re} (1 - a^2 - b^2) d |E(A, a, b) f|^2 \right| \geq \varepsilon \left| \iint_{\Re} d |E(A, a, b) f|^2 \right| \\
&= \varepsilon \left| \iint d |E(A, a, b) f|^2 \right| = \varepsilon \cdot |f|^2. \quad 69)
\end{aligned}$$

Also muß $f=0$ sein, d.h. es ist $E(A, \Re)=0$. — Wenn $\Re, \Im$ zwei punktfremde Rechtecke sind (die Ränder dürfen sich treffen), so ist $E(A, \Re) \cdot E(A, \Im) = 0$ ⁷⁰⁾. Wenn die beiden Rechtecke $\Re, \Im$ zusammen genau ein Rechteck $\mathfrak{T}$ ausmachen, so ist $E(A, \Re) + E(A, \Im) = E(A, \mathfrak{T})$ ⁷¹⁾.

Durch wiederholte Anwendung dieser Tatsachen folgert man nun: wenn $\mathfrak{U}$ eine Summe von endlich vielen Rechtecken ist, so ist erstens

$$E(A, \mathfrak{U}) = E(A, \Re_1) + \dots + E(A, \Re_n)$$

($\Re_1, \dots, \Re_n$ seien die genannten Rechtecke) ein P. O.; zweitens ist er für alle Zerlegungen $\Re_1, \dots, \Re_n$ von $\mathfrak{U}$ derselbe; drittens ist, wenn $\mathfrak{U}, \mathfrak{V}$ zwei solche Mengen sind und höchstens Randpunkte gemein haben, $E(A, \mathfrak{U}) \cdot E(A, \mathfrak{V}) = 0$, $E(A, \mathfrak{U} + \mathfrak{V}) = E(A, \mathfrak{U}) + E(A, \mathfrak{V})$. Außerdem muß, wenn $\mathfrak{U}$ den Einheitskreis nicht trifft, $E(A, \mathfrak{U}) = 0$ sein (weil alle $E(A, \Re_1), \dots, E(A, \Re_n) = 0$ sind); hieraus folgert man mühelos weiter: wenn $\mathfrak{U}, \mathfrak{V}$ dieselben Punkte mit dem Einheitskreis gemein haben (und keine Randecke auf ihm), so ist $E(A, \mathfrak{U}) = E(A, \mathfrak{V})$.

Jetzt sind wir in der Lage, das Stieltjessche Integral $\iint$ sozusagen auf Polarkoordinaten zu transformieren.

⁶⁹⁾ Da f in $\Re(R, a')$ und $\Re(R, a'')$ liegt, ist für $a \leq a'$ oder $\geq a''$ $P_{\Re(R, a)} f = 0$ bzw. $= f$, also $E(A, a, b)f$ von a unabhängig; ebenso ist es für $b \leq b'$ oder $\geq b''$ von b unabhängig. Daher dürfen wir $\iint$ und $\iint_{\Re}$ identifizieren.

⁷⁰⁾ Seien $\Re, \Im$ etwa $a' \leq x \leq a''$, $b' \leq y \leq b''$ und $c' \leq x \leq c''$, $d' \leq y \leq d''$; es muß $a'' \leq c'$ oder $c'' \leq a'$ oder $b'' \leq d'$ oder $d'' \leq b'$ sein. Für $a'' \leq c'$ ist $\Re(R, a'')$ Teil von $\Re(R, c')$, also $\Re(R, a'')$ orth. zu ihm — daher ist $P_{\Re(R, a'')} P_{\Re(R, c')} = 0$ und erst recht $E(A, \Re) \cdot E(A, \Im) = 0$. Ebenso schließt man in den drei anderen Fällen.

⁷¹⁾ Es muß $a'' = c'$ oder $c'' = a'$ und gleichzeitig $b' = d'$, $d'' = b''$ sein, oder $b'' = d'$ oder $d'' = b'$ und gleichzeitig $a' = c'$, $a'' = c''$; man verifiziere an der Definition der $E(A, \Re)$.

6. Es sei $0 \leq \varrho < \sigma \leq 1$, dann gibt es sicher Rechteck-Mengen $\mathfrak{U}$, die mit dem Einheitskreise genau den Bogen $2\pi\varrho, 2\pi\sigma$ gemein haben (und für die keine Randecke auf dem Einheitskreise liegt), und für alle diese ist (nach dem oben Gesagten) $E(A, \mathfrak{U})$ derselbe P. O. — er heiße $F(\varrho, \sigma)$. Es ist (für $0 \leq \varrho < \sigma < \tau \leq 1$) offenbar $F(\varrho, \sigma) + F(\varrho, \tau) = F(\varrho, \tau)$, also, wenn wir $F(\varrho) = F(0, \varrho)$ setzen, $F(\varrho, \sigma) = F(\sigma) - F(\varrho)$. Weiter muß, weil offenbar $F(0, 1) = 1$ ist, $F(1) = 1$ sein. Die $F(\varrho)$ sind P. O., da für $\varrho \leq \sigma$ $F(\sigma) - F(\varrho)$ auch P. O. ist, ist $F(\varrho)$ Teil von $F(\sigma)$. $F(0)$ ist zunächst undefiniert, es sei $= 0$.

Man teile das Integrationsgebiet der Formel

$$(Af, f) = \iint (a + ib) d|E(A, a, b)f|^2$$

(ein großes Rechteck!) in Rechtecke $\mathfrak{R}_1, \dots, \mathfrak{R}_k$ mit den folgenden Eigenschaften ein: Jedes derselben hat einen Durchmesser $\leq \varepsilon$ ($\varepsilon > 0$), keines eine Randecke auf dem Einheitskreise; und zwar seien sie so angeordnet, daß mit dem Einheitskreise $\mathfrak{R}_1, \dots, \mathfrak{R}_k$ ($k \leq \infty$) Punkte gemein haben, und zwar bzw. die Bögen $2\pi\varrho_0, 2\pi\varrho_1; \dots, 2\pi\varrho_{h-1}, 2\pi\varrho_h$ ($0 = \varrho_0 < \varrho_1 < \dots < \varrho_{h-1} < \varrho_h = 1$). Aus jedem $\mathfrak{R}_1, \dots, \mathfrak{R}_k$ wählen wir je einen Punkt $\xi_1, \dots, \xi_k$ aus, dann ist

$$(Af, f) = \iint (a + ib) d|E(A, a, b)f|^2 = \sum_{\nu=1}^k \iint_{\mathfrak{R}_\nu} (a + ib) d|E(A, a, b)f|^2$$

beliebig wenig ⁷²⁾ von

$$\begin{aligned} \sum_{\nu=1}^k \xi_\nu \iint_{\mathfrak{R}_\nu} d|E(A, a, b)f|^2 &= \sum_{\nu=1}^k \xi_\nu \iint_{\mathfrak{R}_\nu} d(E(A, a, b)f, f) = \sum_{\nu=1}^k \xi_\nu (E(A, \mathfrak{R}_\nu)f, f) \\ &= \sum_{\nu=1}^h \xi_\nu \cdot |F(\varrho_{\nu-1}, \varrho_\nu)f|^2 = \sum_{\nu=1}^h \xi_\nu \cdot (|F(\varrho_\nu)f|^2 - |F(\varrho_{\nu-1})f|^2) \end{aligned}$$

verschieden. Wir können natürlich $\xi_\nu = e^{2\pi i \cdot \varrho'_\nu}$, $\varrho_{\nu-1} < \varrho'_\nu < \varrho_\nu$ ($\nu = 1, \dots, h$) wählen, dann ist unser Ausdruck

$$\sum_{\nu=1}^h e^{2\pi i \cdot \varrho'_\nu} (|F(\varrho_\nu)f|^2 - |F(\varrho_{\nu-1})f|^2).$$

Dabei sind die $\varrho_0, \varrho_1, \dots, \varrho_{h-1}, \varrho_h$ voneinander $\leq \varepsilon$ entfernt.

Hieraus folgt aber

$$(Af, f) = \int_0^1 e^{2\pi i \cdot \varrho} d|F(\varrho)f|^2. \quad 73)$$

⁷²⁾ Da $\oint d|E(A, a, b)f|^2$ für alle Rechtecke $\mathfrak{S} \geq 0$ ist (nämlich $= |E(A, \mathfrak{S})f|^2$),

höchstens $\varepsilon \iint |E(A, a, b)f|^2 = \varepsilon \cdot |f|^2$.

⁷³⁾ Das Integral rechts konvergiert, z. B. weil wir $\varrho_0, \varrho_1, \dots, \varrho_{h-1}, \varrho_h$ vorschreiben und die Einteilung $\mathfrak{R}_1, \dots, \mathfrak{R}_k$ daran anpassen können.

Satz 12*. Zu einem unitären Operator U gibt es eine Schar von P.O. $E(\lambda)$ ($0 \leq \lambda \leq 1$) mit den folgenden Eigenschaften:

α) Wenn $\lambda \leq \mu$ ist, so ist $E(\lambda)$ Teil von $E(\mu)$; $E(0) = 0$, $E(1) = 1$.

β) Wenn $\lambda \geq \lambda_0$, $\lambda \rightarrow \lambda_0$, so gilt für jedes f $E(\lambda)f \rightarrow E(\lambda_0)f$.

Dabei gilt für alle f, g

$$(Uf, g) = \int_0^1 e^{2\pi i \cdot \lambda} d(E(\lambda)f, g).$$

Beweis. Betrachten wir die vorhin konstruierten $F(\lambda)$. Sie haben die Eigenschaft α), aber β) eventuell nicht. Sei $0 \leq \lambda_0 < 1$ und $\lambda_1 > \lambda_2 > \dots > \lambda_0$, $\lambda_n \rightarrow \lambda$, dann bilden die $F(\lambda_1), F(\lambda_2), \dots$ eine abnehmende Folge von P.O., nach Satz 19 gibt es also einen P.O. E , so daß stets $E(\lambda_n)f \rightarrow Ef$ ist; E hängt nur von λ_0 ab, denn zu zwei Folgen $\lambda_1, \lambda_2, \dots$ muß dasselbe E gehören, weil geeignete Teilfolgen von ihnen zu einer ebensolchen Folge zusammengefaßt werden können — für die im obigen Sinne die Konvergenz auch stattfinden muß. Dieses E heiße $E(\lambda_0)$.

Da $F(\lambda_0)$ Teil aller $F(\lambda_1), F(\lambda_2), \dots$ ist, ist es auch Teil von $E(\lambda_0)$; wenn $\lambda_0 < \mu_0$ ist, so können wir $\lambda_1 = \mu_0$ wählen, so daß $E(\lambda_0)$ Teil von $F(\mu_0) = F(\lambda_1)$ ist (weil es alle $F(\lambda_1), F(\lambda_2), \dots$ sind). Also: $F(\lambda)$ Teil von $E(\lambda)$, für $\lambda < \mu$ $E(\lambda)$ Teil von $F(\mu)$, daher für $\lambda \leq \mu$ $E(\lambda)$ Teil von $E(\mu)$.

Sei nun $\lambda > \lambda_0$, $\lambda \rightarrow \lambda_0$. $E(\lambda) - E(\lambda_0)$ ist Teil von $F(\lambda) - E(\lambda_0)$, also $|E(\lambda)f - E(\lambda_0)f| \leq |F(\lambda)f - E(\lambda_0)f|$, und da letzteres (nach Definition von $E(\lambda_0)$) gegen 0 strebt, muß $E(\lambda)f \rightarrow E(\lambda_0)f$. Die $E(\lambda)$ erfüllen also β) und auch α) bis auf $E(0) = 0$, $E(1) = 1$. Wir erzwingen dies durch Definition — dann geht β) für $\lambda = 0$ eventuell verloren — das alte $E(0)$ heiße $\bar{E}$.

Für $F(\lambda)$ hatten wir die Relation

$$(Uf, f) = \int_0^1 e^{2\pi i \cdot \lambda} d|F(\lambda)f|^2.$$

Da $|F(\lambda)f|^2$, $|E(\lambda)f|^2$ monoton sind (in λ) und für $\lambda < \mu$ $|F(\lambda)f|^2 \leq |E(\lambda)f|^2 \leq |F(\mu)f|^2$ gilt, überlegt man sich leicht, daß sie auch für $E(\lambda)$ besteht. Dies ist aber die Schlußrelation des Satzes für $f = g$. Wenn man $\frac{f+g}{2}$ und $\frac{f-g}{2}$ einsetzt und subtrahiert, so erhält man ihren Realteil für beliebige f, g , ersetzt man f, g durch if, g , dann entsteht der Imaginärteil.

Es bleibt noch die Verletzung β) für $\lambda = 0$ zu beheben, dies gelingt, indem man für $0 < \lambda < 1$ $E(\lambda)$ durch $E(\lambda) - \bar{E}$ ersetzt — dann bleibt alles andere beim alten und dieser Punkt kommt in Ordnung.

7. Der Satz 12* war die Hauptschwierigkeit, nun folgen einige leichte Eindeutigkeitssätze.

Satz 13*. Wenn die P. O. $E(\lambda)$ den Bedingungen $\alpha), \beta)$ von Satz 12* entsprechend gewählt werden, so gibt es genau einen unitären Operator U , der zu ihnen im Verhältnis von Satz 12* steht.

Beweis. Daß es höchstens ein solches U gibt, ist klar: Satz 12* legt alle (Uf, g) , also alle Uf fest. Bleibt die Existenz zu beweisen.

Durch dieselben Überlegungen, wie sie beim Beweise von Satz 36 angestellt wurden, gelingt (wegen $|e^{2\pi i \cdot \lambda}| = 1$) die Abschätzung

$$\left| \int_0^1 e^{2\pi i \cdot \lambda} d(E(\lambda) f, g) \right| \leq |f| |g|.$$

Nennen wir das Integral links $L(g)$ (f fest, g variabel), so ist

$$L(a_1 g_1 + \dots + a_n g_n) = \bar{a}_1 L(g_1) + \dots + \bar{a}_n L(g_n),$$

$$|L(g)| \leq |f| |g| = C \cdot |g|,$$

der F. Rieszsche Satz ist anwendbar (vgl. Anm. 53): Es gibt genau ein f^* mit $L(g) = (f^*, g)$. Wir definieren U durch $Uf = f^*$, es ist offenbar ein überall sinnvoller lin. O. — wir müssen nur noch die Unitarität von U beweisen.

Zunächst ist $|(Uf, g)| \leq |f| |g|$, also U beschränkt; wir bilden U^* , und haben:

$$(Uf, g) = \int_0^1 e^{2\pi i \lambda} d(E(\lambda) f, g), \quad (U^* f, g) = \int_0^1 e^{-2\pi i \lambda} d(E(\lambda) f, g)$$

(die erste Gleichung ist Definition, die zweite entsteht durch Vertauschen von f, g und Nehmen der komplex-Konjugierten). Genau wie am Ende von § 4 berechnet man auf Grund dieser Formeln $(Uf, Ug) = (f, g)$, $(U^* f, U^* g) = (f, g)$. Daraus folgt $(U^* U f, g) = (U U^* f, g) = (f, g)$ für alle f, g , also $U^* U = U U^* = 1$ — also die Unitarität.

Satz 14*. Zu einem gegebenen unitären Operator U kann nur eine Schar $E(\lambda)$ von P. O. im Verhältnisse des Satzes 12* stehen.

Beweis. Die Formel

$$[(U^n f, g) = \int_0^1 e^{2\pi i n \lambda} d(E(\lambda) f, g)$$

ist für $n = 0$ trivial, für $n = 1$ Definition, für $n = -1$ Folge von $U^{-1} = U^*$. Wenn sie für m, n bewiesen ist, so gilt sie auch für $m - n$: denn mit der Methode vom Schlusse von § 4 ergibt sich

$$(U^m f, U^n g) = \int_0^1 e^{2\pi i (m-n) \lambda} d(E(\lambda) f, g),$$

und es ist

$$(U^m f, U^n g) = ((U^n)^* U^m f, g) = (U^{*n} U^m f, g) = (U^{-n} U^m f, g) = (U^{m-n} f, g).$$

Daher gilt sie für alle $n = 0, \pm 1, \pm 2, \dots$

Wenn $F(\lambda)$ eine zweite Schar von P.O. mit denselben Eigenschaften ist, so gilt die obige Gleichung auch mit $F(\lambda)$ an Stelle von $E(\lambda)$. Also gilt

$$\int_0^1 p(\lambda) d((E(\lambda) f, g) - (F(\lambda) f, g)) = 0$$

zunächst für alle $p(\lambda) = e^{2\pi i \lambda}$, und dann auch für ihre Linearaggregate, die trigonometrischen Polynome. Aus Stetigkeitsgründen gilt es aber auch für alle stetigen Funktionen $p(\lambda)$ mit $p(0) = p(1)$. Wegen der Bedingung β) in Satz 12* gilt sie sogar noch für solche $p(\lambda)$, die endlich viele Sprünge aufweisen — vorausgesetzt, daß die Unstetigkeit stets rechts von der Sprungstelle liegt. Wir setzen nun

$$p(\lambda) = \begin{cases} 1, & \text{für } 0 < \lambda \leq \lambda_0, \\ 0, & \text{sonst} \end{cases},$$

($0 \leq \lambda_0 < 1$), dann haben wir

$$\int_0^{\lambda_0} d((E(\lambda) f, g) - (F(\lambda) f, g)) = (E(\lambda_0) f, g) - (F(\lambda_0) f, g) = 0;$$

da dies für alle f, g gilt, ist somit $E(\lambda_0) = F(\lambda_0)$. (Dabei muß $\lambda_0 \neq 1$ sein, für $\lambda_0 = 1$ ist es aber wegen α) in Satz 12* der Fall.)

8. Wir haben durch Satz 12*, 13*, 14* gezeigt: die von Satz 12* gegebene Zuordnung eines unitären Operators U zu einer Schar von P.O. $E(\lambda)$ mit den dortigen Eigenschaften α), β) ist ein-eindeutig, und alle U und alle $E(\lambda)$ werden dadurch erschöpft. Zur Z. d. E. fehlt diesen $E(\lambda)$ noch die Eigenschaft $E(\lambda) f \rightarrow f$ für $\lambda < 1, \lambda \rightarrow 1$ (vgl. Def. 17), wir haben daher alles in Kap. IX Angekündigte bewiesen, wenn wir noch zeigen, daß gerade diese Eigenschaft damit gleichbedeutend ist, daß $\varphi = U\varphi$ nur für $\varphi = 0$ stattfindet.

Sei zuerst $\varphi = U\varphi, \varphi \neq 0$. Dann ist

$$\begin{aligned} \Re((\varphi, \varphi) - (U\varphi, \varphi)) &= \Re \int_0^1 (1 - e^{2\pi i \lambda}) d(E(\lambda) \varphi, \varphi) \\ &= \int_0^1 2 \sin^2 \pi \lambda d|E(\lambda) \varphi|^2. \end{aligned}$$

Dies muß verschwinden, und ist

$$\begin{aligned} &\geq \int_{\varepsilon}^{1-\varepsilon} 2 \sin^2 \pi \lambda d|E(\lambda) \varphi|^2 \geq 2 \sin^2 \pi \varepsilon \int_{\varepsilon}^{1-\varepsilon} d|E(\lambda) \varphi|^2 \\ &= 2 \sin^2 \pi \varepsilon (|E(1-\varepsilon) \varphi|^2 - |E(\varepsilon) \varphi|^2) \end{aligned}$$

(für alle $\varepsilon > 0$). Für $\varepsilon \leq \frac{1}{2}$ ist die rechte Seite gewiß ≥ 0 , also $= 0$:

$$|E(1 - \varepsilon)\varphi| = |E(\varepsilon)\varphi|.$$

Da die rechte Seite mit $\varepsilon \rightarrow 0$ gegen 0 strebt, gilt dasselbe von der linken: $E(1 - \varepsilon)\varphi \rightarrow 0 \neq \varphi$.

Sei umgekehrt für $\lambda < 1$, $\lambda \rightarrow 1$ nicht $E(\lambda)f \rightarrow f$. Nach Satz 19 existiert ein P.O. E , so daß für alle g $E(\lambda)g \rightarrow Eg$, es ist also $Ef \neq f$. Für $f' = f - Ef$ haben wir daher: $f' \neq 0$, $E(\lambda)f' \rightarrow Ef' = 0$. Nun ist $|E(\lambda)f'|$ nie negativ, mit λ monoton nicht fallend — da es für $\lambda \rightarrow 1$ gegen 0 strebt, muß es stets $= 0$ sein. Also ist $E(\lambda)f' = 0$ für alle $\lambda < 1$, und natürlich $E(1)f' = f'$. Hieraus folgt

$$(Uf', g) = \int_0^1 e^{2\pi i \lambda} d(E(\lambda)f', g) = (f', g),$$

dies gilt für alle g , also ist $Uf' = f'$.

Damit ist alles bewiesen.

Anhang III. Operatoren und Matrizen.

1. Es ist wohl erwünscht, unsere Resultate über H.O. in die übliche Terminologie der Matrizen zu übersetzen. Es wird sich zeigen, daß im Unbeschränkten die der Operatoren bedeutend praktischer ist. (Im Beschränkten kommt beides, nach Hilberts Resultaten, aufs selbe heraus.) Im übrigen werden wir erkennen, daß die abg. lin. H. O. und die Hermiteschen Matrizen mit lauter endlichen (aber nicht notwendig beschränkten!) Zeilen-Absolutwertquadratsummen — die sogenannten „quadrierbaren Matrizen“ — einander zugeordnet werden können.

Die genauere Theorie dieser Matrizen soll in der in Anm.²²⁾ angekündigten Arbeit ausgeführt werden, hier soll nur soviel gegeben werden, als zur Orientierung nötig ist, und was wir für die Zwecke von Anhang IV brauchen.

Eine Matrix $A = \{a_{\mu\nu}\}$ ($\mu, \nu = 1, 2, \dots$) heiße Hermitesch quadrierbar (kurz: H. quadr.), wenn $a_{\mu\nu} = \overline{a_{\nu\mu}}$ gilt, und alle $\sum_{\nu=1}^{\infty} |a_{\mu\nu}|^2$ ($\mu = 1, 2, \dots$) endlich sind⁷⁴⁾. Wenn $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System ist, so bilden wir einen Operator R , der für die φ_μ (und nur diese) Sinn hat, und zwar $R\varphi_\mu = \sum_{\nu=1}^{\infty} a_{\mu\nu}\varphi_\nu$ (die Reihe konvergiert). Die von den $\varphi_1, \varphi_2, \dots$ aufgespannte abg. lin. M. ist $\mathfrak{H}$, und es ist

$$(\varphi_\mu, R\varphi_\nu) = \overline{a_{\nu\mu}} = a_{\mu\nu} = (R\varphi_\mu, \varphi_\nu),$$

⁷⁴⁾ Dann sind auch alle $\sum_{\nu=1}^{\infty} a_{\mu\nu} \overline{a_{\nu\mu}}$ absolut konvergent, d. h. die Matrix A^2 kann gebildet werden.

d. h. R ist ein H. O. Wir nennen R den elementaren H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$; und $\hat{R}, \tilde{R}$ den lin. bzw. abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$. — $\tilde{R}f$ hat offenbar für die $\sum_{\mu=1}^n x_\mu \varphi_\mu$ Sinn (und nur für diese), und es ist

$$\tilde{R}f = \sum_{\mu=1}^n x_\mu \left(\sum_{\nu=1}^{\infty} a_{\mu\nu} \varphi_\nu \right) = \sum_{\nu=1}^{\infty} \left[\sum_{\mu=1}^n a_{\mu\nu} x_\mu \right] \varphi_\nu.$$

$\tilde{R}$ ist nicht so leicht direkt zu charakterisieren, wohl aber seine Erweiterungselemente. Damit f eines sei, und zwar das f^* ihm zugeordnet, muß

$$(f^*, \varphi_\nu) = (f, R\varphi_\nu) = \sum_{\mu=1}^{\infty} (f, \varphi_\mu) (\varphi_\mu, R\varphi_\nu) = \sum_{\mu=1}^{\infty} a_{\mu\nu} (f, \varphi_\mu) \quad (\mu = 1, 2, \dots)$$

sein (R und $\tilde{R}$ haben ja dieselben Erweiterungselemente). Wir schreiben

$$(f, \varphi_\mu) = x_\mu \quad (\mu = 1, 2, \dots), \quad f = \sum_{\mu=1}^{\infty} x_\mu \varphi_\mu,$$

dann existiert ein f^* mit $(f^*, \varphi_\nu) = \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu$ offenbar dann und nur dann, wenn $\sum_{\nu=1}^{\infty} \left| \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \right|^2$ endlich ist, und zwar ist es dann gleich $\sum_{\nu=1}^{\infty} \left(\sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \right) \varphi_\nu$. Wir haben also: $f = \sum_{\mu=1}^{\infty} x_\mu \varphi_\mu$ ($\sum_{\mu=1}^{\infty} |x_\mu|^2$ endlich) ist dann und nur dann Erweiterungselement, wenn für die

$$y_\nu = \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \quad ^{75})$$

$\sum_{\mu=1}^{\infty} |y_\mu|^2$ endlich ist, und zwar ist dann $f^* = \sum_{\mu=1}^{\infty} y_\mu \varphi_\mu$. (Wenn sogar $\tilde{R}f$ Sinn hat, ist es erst recht dieses.)

Weiter ist

$$(f^*, f) = \sum_{\nu=1}^{\infty} y_\nu \bar{x}_\nu = \sum_{\nu=1}^{\infty} \left(\sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \right) \bar{x}_\nu = \sum_{\nu=1}^{\infty} \left(\sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \bar{x}_\nu \right),$$

$$(f, f^*) = \sum_{\mu=1}^{\infty} x_\mu \bar{y}_\mu = \sum_{\mu=1}^{\infty} x_\mu \left(\sum_{\nu=1}^{\infty} \bar{a}_{\nu\mu} \bar{x}_\nu \right) = \sum_{\mu=1}^{\infty} \left(\sum_{\nu=1}^{\infty} a_{\mu\nu} x_\mu \bar{x}_\nu \right).$$

Die 0-Klasse ist also durch die Vertauschbarkeit von $\sum_{\mu=1}^{\infty}, \sum_{\nu=1}^{\infty}$ gekennzeichnet!

2. Wir zeigen nun, daß zu jedem abg. lin. H. O. R eine H. quadr. Matrix $A = \{a_{\mu\nu}\}$ und ein vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$ existiert, so daß R der abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ ist. Es genügt $\varphi_1, \varphi_2, \dots$ anzugeben, A ist ja durch $a_{\mu\nu} = (R\varphi_\mu, \varphi_\nu)$ festgelegt, und wenn die $R\varphi_1, R\varphi_2, \dots$ alle Sinn haben, ist A H. quadr.: $a_{\mu\nu} = \overline{a_{\nu\mu}}$ ist klar, und $\sum_{\nu=1}^{\infty} |a_{\mu\nu}|^2 = \sum_{\nu=1}^{\infty} |(R\varphi_\mu, \varphi_\nu)|^2 = |R\varphi_\mu|^2$ ist endlich.

⁷⁵⁾ $\sum_{\mu=1}^{\infty} |a_{\mu\nu}|^2, \sum_{\mu=1}^{\infty} |x_\mu|^2$ sind endlich, also konvergieren die $\sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu$ jedenfalls absolut!

Sei S der elementare Operator von $A, \{\varphi_1, \varphi_2, \dots\}$, dann ist

$$S\varphi_\mu = \sum_{\nu=1}^{\infty} a_{\mu\nu} \varphi_\nu = \sum_{\nu=1}^{\infty} (R\varphi_\mu, \varphi_\nu) \cdot \varphi_\nu = R\varphi_\mu;$$

d. h. S ist das auf die $\varphi_1, \varphi_2, \dots$ eingeschränkte R . Das, was wir durch geeignete Wahl der $\varphi_1, \varphi_2, \dots$ erreichen sollen, ist also $\tilde{S} = R$, d. h.: Daß zu jedem f mit sinnvollem Rf eine Folge von Linearaggregaten je endlichvieler $\varphi_1, \varphi_2, \dots$ existiere, $f_1, f_2, \dots$, so daß $f_n \rightarrow f, Rf_n \rightarrow Rf$ (natürlich müssen alle $R\varphi_n$ sinnvoll sein). Dies ist erreicht, wenn wir eine Folge $g_1, g_2, \dots$ mit lauter sinnvollen $Rg_1, Rg_2, \dots$ haben, derart, daß zu jedem f mit sinnvollem Rf eine Teilfolge $g_{N_1}, g_{N_2}, \dots$ mit $g_{N_n} \rightarrow f, Rg_{N_n} \rightarrow Rf$ existiert (diese ist dann von selbst überall dicht): Dann genügt es, $f_1, f_2, \dots$ nach Satz 8 zu $\varphi_1, \varphi_2, \dots$ zu „orthogonalisieren“.

Eine solche Folge $g_1, g_2, \dots$ finden wir so: Sei $h_1, h_2, \dots$ überall dicht, aber sonst beliebig. Zu jedem Zahlentripel m, n, k ($= 1, 2, \dots$) zu dem es ein f mit sinnvollem Rf und $|h_m - f| \leq \frac{1}{k}, |h_n - Rf| \leq \frac{1}{k}$ gibt, ordnen wir ein solches f zu: $f_{m,n,k}$. Diese $f_{m,n,k}$ leisten (als Folge geschrieben) alles Gewünschte. Denn sei f, Rf und $\varepsilon > 0$ gegeben. Wir wählen ein k mit $\frac{2}{k} \leq \varepsilon$ und h_m und h_n mit $|h_n - f| \leq \frac{1}{k}, |h_m - Rf| \leq \frac{1}{k}$. Dann ist $f_{m,n,k}$ gewiß definiert, und $|h_m - f_{m,n,k}| \leq \frac{1}{k}, |h_n - Rf_{m,n,k}| \leq \frac{1}{k}$, also $|f - f_{m,n,k}| \leq \varepsilon, |Rf - Rf_{m,n,k}| \leq \varepsilon$. Da $\varepsilon > 0$ beliebig ist, sind wir am Ziele.

3. Zu gegebenem R kann es auch mehrere $A, \{\varphi_1, \varphi_2, \dots\}$ geben, deren abg. lin. H. O. dieses R ist. Seien etwa $A = \{a_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$ und $B = \{b_{\mu\nu}\}, \{\psi_1, \psi_2, \dots\}$ derart. (Sowohl A wie B sollen H. quadr. sein!) Die beiden vollst. norm. orth. Systeme $\varphi_1, \varphi_2, \dots, \psi_1, \psi_2, \dots$ bestimmen eine unitäre unendliche Matrix $\{u_{\mu\nu}\}$ mit $u_{\mu\nu} = (\varphi_\mu, \psi_\nu)$. In der Tat ist:

$$(U_1) \quad \sum_{\varrho=1}^{\infty} u_{\mu\varrho} \overline{u_{\nu\varrho}} = \sum_{\varrho=1}^{\infty} (\varphi_\mu, \psi_\varrho) (\psi_\varrho, \varphi_\nu) = (\varphi_\mu, \varphi_\nu) = \begin{cases} 1, & \text{für } \mu = \nu, \\ 0, & \text{für } \mu \neq \nu, \end{cases}$$

$$\sum_{\varrho=1}^{\infty} u_{\varrho\mu} \overline{u_{\varrho\nu}} = \sum_{\varrho=1}^{\infty} (\psi_\varrho, \varphi_\mu) (\varphi_\varrho, \psi_\nu) = (\psi_\nu, \psi_\mu) = \begin{cases} 1, & \text{für } \mu = \nu, \\ 0, & \text{für } \mu \neq \nu. \end{cases}$$

Und weiter:

$$(U_2) \quad \varphi_\mu = \sum_{\nu=1}^{\infty} (\varphi_\mu, \psi_\nu) \cdot \psi_\nu = \sum_{\nu=1}^{\infty} u_{\mu\nu} \cdot \psi_\nu, \quad \psi_\mu = \sum_{\nu=1}^{\infty} (\psi_\mu, \varphi_\nu) \cdot \varphi_\nu = \sum_{\nu=1}^{\infty} \overline{u_{\nu\mu}} \cdot \varphi_\nu.$$

Auch für $a_{\mu\nu}, b_{\mu\nu}$ gelten die gewohnten Transformationsformeln:

$$a_{\mu\nu} = (\varphi_\mu, R\varphi_\nu) = \sum_{\varrho=1}^{\infty} (\varphi_\mu, \psi_\varrho) (\psi_\varrho, R\varphi_\nu) = \sum_{\varrho=1}^{\infty} (\varphi_\mu, \psi_\varrho) (R\psi_\varrho, \varphi_\nu)$$

$$= \sum_{\varrho=1}^{\infty} (\varphi_\mu, \psi_\varrho) \left[\sum_{\sigma=1}^{\infty} (R\psi_\varrho, \psi_\sigma) (\psi_\sigma, \varphi_\nu) \right] = \sum_{\varrho=1}^{\infty} \left(\sum_{\sigma=1}^{\infty} b_{\varrho\sigma} u_{\mu\varrho} \overline{u_{\nu\sigma}} \right).$$

Wenn wir μ, ν (und ϱ, σ) vertauschen und die komplex-Konjugierten nehmen, so wird (wegen $a_{\mu\nu} = \overline{a_{\nu\mu}}$, $b_{\varrho\sigma} = \overline{b_{\sigma\varrho}}$)

$$(U_3) \quad a_{\mu\nu} = \sum_{\varrho=1}^{\infty} \left(\sum_{\sigma=1}^{\infty} b_{\varrho\sigma} u_{\mu\varrho} \overline{u_{\nu\sigma}} \right) = \sum_{\sigma=1}^{\infty} \left(\sum_{\varrho=1}^{\infty} b_{\varrho\sigma} u_{\varrho\sigma} \overline{u_{\nu\sigma}} \right).$$

Ebenso erhalten wir die Umkehrungen:

$$(U_4) \quad b_{\mu\nu} = \sum_{\varrho=1}^{\infty} \left(\sum_{\sigma=1}^{\infty} a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right) = \sum_{\sigma=1}^{\infty} \left(\sum_{\varrho=1}^{\infty} b_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right).$$

Bis hierher gilt also die gewöhnliche unitäre Transformationstheorie der Hermiteschen Matrizen, aber nicht weiter! Denn wenn wir eine H. quadr. Matrix $A = \{a_{\mu\nu}\}$ und ein vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$ haben, und irgendeine unitäre Transformationsmatrix $\{u_{\mu\nu}\}$ (nach (U_1)) darauf anzuwenden versuchen (nach (U_2) , (U_3) , (U_4)), so braucht zunächst nichts zu konvergieren. Aber auch wenn die in (U_3) , (U_4) angegebenen Konvergenzen alle stattfinden, so kann sich der abg. lin. H. O. von $A_1\{\varphi_1, \varphi_2, \dots\}$ bei der Transformation ändern. Wir gehen auf diese Verhältnisse hier nicht näher ein, da sie an anderem Orte (vgl. Anm. ²²) genau diskutiert werden sollen.

4. Es sei noch erwähnt, wie sich die abg. lin. M. $\mathfrak{E}, \mathfrak{F}$ die zum abg. lin. H. O. $\tilde{R}$ von $A, \{\varphi_1, \varphi_2, \dots\}$ gehören (R sei sein elementarer H. O.) bestimmen lassen (vgl. Satz 22). $\mathfrak{E}$ bzw. $\mathfrak{F}$ ist die Menge aller $(\tilde{R} \pm i1)f$. Aus der abg. von $\tilde{R}$ und den in Satz 22 auseinandergesetzten Eigenschaften von $\tilde{R} \pm i1$ folgt man leicht, daß dies die Menge der Häufungspunkte von $(\tilde{R} \pm i1)f$ ist⁷⁶). Dies aber ist die von den $(R \pm i1)f$, d. h. von den $(R \pm i1)\varphi_\mu$, aufgespannte lin. M.; also ist $\mathfrak{E}$ bzw. $\mathfrak{F}$ die von den $(R \pm i1)\varphi_\mu$ aufgespannte abg. lin. M., d. h. von den $\sum_{\nu=1}^{\infty} a_{\mu\nu} \varphi_\nu \pm i \varphi_\mu$.

Anhang IV. Der Operator $\bar{R}$.

1. $\bar{R}$ und seine Cayleysche Transformierte $\bar{U}$ wurden in Kap. X (Satz 37) folgendermaßen definiert: Sei $\varphi_0, \varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System, dann ist

$$\bar{U} \left(\sum_{\nu=0}^{\infty} x_\nu \varphi_\nu \right) = \sum_{\nu=0}^{\infty} x_\nu \varphi_{\nu+1} \quad \left(\sum_{\nu=1}^{\infty} |x_\nu|^2 \text{ endlich} \right);$$

d. h. $\bar{R}f$ hat Sinn für alle $f = \varphi - \bar{U}\varphi = x_0\varphi_0 + (x_1 - x_0)\varphi_1 + (x_2 - x_1)\varphi_2 + \dots$, und dann ist $\bar{R}f = i(\varphi + \bar{U}\varphi) = ix_0\varphi_0 + i(x_0 + x_1)\varphi_1 + i(x_1 + x_2)\varphi_2 + \dots$.

Oder auch: $\bar{R} \left(\sum_{\nu=0}^{\infty} y_\nu \varphi_\nu \right)$ hat Sinn, wenn $|y_0|^2 + |y_0 + y_1|^2 + |y_0 + y_1 + y_2|^2 + \dots$

⁷⁶) Denn wenn eine Folge $(\tilde{R} \pm i1)f$ konvergiert, so konvergiert auch $U^{\pm 1}(\tilde{R} \pm i1)f = (\tilde{R} \mp i1)f$, also f und $\tilde{R}f$ — und natürlich auch umgekehrt: Wenn $f, \tilde{R}f$ konvergieren, konvergiert $(\tilde{R} \pm i1)f$.

konvergiert (es muß ja $x_0 = y_0$, $x_1 = y_0 + y_1$, $x_2 = y_0 + y_1 + y_2$, ... sein), und zwar ist es dann gleich $i y_0 \varphi_0 + (2i y_0 + i y_1) \varphi_1 + (2i y_0 + 2i y_1 + i y_2) \varphi_2 + \dots$. Wir haben hier eine Art Matrix von $\bar{R}$ gewonnen:

$$a_{\mu\nu} = \begin{cases} 0, & \text{für } \mu > \nu, \\ i, & \text{für } \mu = \nu, \\ 2i, & \text{für } \mu < \nu. \end{cases}$$

Freilich ist sie es nicht im Sinne von Anhang III: Sie ist weder H. noch quadr., was kein Wunder ist, da die $\bar{R} \varphi_\mu$, $\mu = 0, 1, 2, \dots$, alle sinnlos sind.

Um eine richtige Matrix von $\bar{R}$ angeben zu können, müssen wir vor allem ein vollst. norm. orth. System $\psi_1, \psi_2, \dots$ finden, für welches alle $R \psi_\mu$, $\mu = 1, 2, \dots$, Sinn haben.

Wir setzen:

$$\omega_{m,k} = \varphi_{m \cdot 2^{k+1}} + \varphi_{m \cdot 2^{k+1} + 1} + \dots + \varphi_{(2m+1) \cdot 2^k - 1} - \varphi_{(2m+1) \cdot 2^k} \\ - \varphi_{(2m+1) \cdot 2^k + 1} - \dots - \varphi_{(m+1) \cdot 2^{k+1} - 1} \quad (m, k = 0, 1, 2, \dots).$$

Nach unserem Kriterium ist $\bar{R} \left(\sum_{\nu=0}^{\infty} x_\nu \varphi_\nu \right)$ gewiß sinnvoll, wenn nur endlich viele $x_\nu \neq 0$ sind und $\sum_{\nu=1}^{\infty} x_\nu = 0$ ist — dies ist aber hier der Fall.

$(\omega_{m,k}, \omega_{n,l}) = 0$ gilt sicher, wenn $\omega_{m,k}, \omega_{n,l}$ kein gemeinsames φ_ν -Glieder haben, und auch wenn alle φ_ν des einen im anderen vorkommen, und zwar alle mit dem Koeffizienten $+1$, oder alle mit -1 . Dies ist aber immer der Fall, außer für $m = n, k = l$, dagegen ist $|\omega_{m,k}|^2$ offenbar $= 2^{k+1}$. Die $\psi_{m,k} = 2^{-\frac{1}{2}(k+1)} \omega_{m,k}$ bilden also ein norm. orth. System, und alle $\bar{R} \psi_{m,k}$ sind sinnvoll. Aber dieses System ist auch vollst., denn sei $f = \sum_{\nu=0}^{\infty} x_\nu \varphi_\nu$ zu allen $\psi_{m,k}$, d. h. zu allen $\omega_{m,k}$ orth. Das bedeutet:

$$x_{m \cdot 2^{k+1}} + x_{m \cdot 2^{k+1} + 1} + \dots + x_{(2m+1) \cdot 2^k - 1} - x_{(2m+1) \cdot 2^k} - x_{(2m+1) \cdot 2^k + 1} - \dots \\ - x_{(m+1) \cdot 2^{k+1} - 1} = 0.$$

Also! $x_0 - x_1 = 0$, $x_2 - x_3 = 0$, $x_4 - x_5 = 0$, $x_6 - x_7 = 0$, ..., d. h. $x_0 = x_1$, $x_2 = x_3$, $x_4 = x_5$, $x_6 = x_7$, ... Weiter: $x_0 + x_1 - x_2 - x_3 = 0$, $x_4 + x_5 - x_6 - x_7 = 0$, ..., d. h. $2x_0 = 2x_2$, $2x_4 = 2x_6$, ..., d. h. $x_0 = x_2 = x_4 = x_6$, $x_4 = x_6 = x_8 = x_{10}$, ... Weiter: $x_0 + x_1 + x_2 + x_3 - x_4 - x_5 - x_6 - x_7 = 0$, d. h. $4x_0 = 4x_4$, ..., d. h. $x_0 = x_4 = \dots = x_8, \dots$. Somit gilt überhaupt $x_0 = x_1 = x_2 = \dots$, und $\sum_{\nu=0}^{\infty} |x_\nu|^2$ kann nur endlich sein, wenn alle $= 0$ sind: also $f = 0$. Damit ist die Vollst. bewiesen.

2. Wenn wir nun die Matrix $\{b_{mk, nl}\}$ durch

$$b_{mk, nl} = (\psi_{m,k}, \bar{R} \psi_{n,l}) = 2^{-\frac{1}{2}(k+l)-1} (\omega_{m,k}, \bar{R} \omega_{n,l})$$

definieren, so ist sie jedenfalls H. quadr. und ihr elementarer H. O., S , ist das auf die $\psi_{m,k}$ eingeschränkte $\bar{R}$. Also ist $\bar{R}$ Forts. von S , und (weil es abg. lin. ist) von $\tilde{S}$. Wenn $\tilde{S}$ max. ist, so muß $\bar{R} = \tilde{S}$ sein, d. h. der abg. lin. H. O. von $\{b_{m,k/n,l}\}, \{\psi_{m,k}\}$.

Wir beweisen dies, indem wir zeigen: das $\mathfrak{E}$ von $\tilde{S}$ ist $= \mathfrak{S}$, d. h. (vgl § 4 von Anhang III) die $\bar{R}\psi_{m,k} + i\psi_{m,k}$, oder auch die $\bar{R}\omega_{m,k} + i\omega_{m,k}$, spannen die abg. lin. M. $\mathfrak{S}$ auf. Oder: wenn $f = \sum_{v=0}^{\infty} x_v \varphi_v$ zu allen $\bar{R}\omega_{m,k} + i\omega_{m,k}$ orth. ist, so ist $f = 0$. Nach § 1 können wir ja ausrechnen:

$$\begin{aligned}\bar{R}\omega_{m,k} &= i\varphi_{m \cdot 2^{k+1}} + 3i\varphi_{m \cdot 2^{k+1}+1} + \dots + (2^{k+1} - 1)i\varphi_{(2m+1) \cdot 2^k - 1} \\ &\quad + (2^{k+1} - 1)i\varphi_{(2m+1) \cdot 2^k} + (2^{k+1} - 3)i\varphi_{(2m+1) \cdot 2^k + 1} + \dots \\ &\quad + 3i\varphi_{(m+1) \cdot 2^{k+1} - 2} + i\varphi_{(m+1) \cdot 2^{k+1} - 1}, \\ \bar{R}\omega_{m,k} + i\omega_{m,k} &= 2i\varphi_{m \cdot 2^{k+1}} + 4i\varphi_{m \cdot 2^{k+1}+1} + \dots + 2^{k+1}i\varphi_{(2m+1) \cdot 2^k - 1} \\ &\quad + (2^{k+1} - 2)i\varphi_{(2m+1) \cdot 2^k} + (2^{k+1} - 4)i\varphi_{(2m+1) \cdot 2^k + 1} + \dots \\ &\quad + 2i\varphi_{(m+1) \cdot 2^{k+1} - 2}.\end{aligned}$$

Daß f zu diesen orth. ist, bedeutet

$$2ix_{m \cdot 2^{k+1}} + 4ix_{m \cdot 2^{k+1}+1} + \dots + 2^{k+1}ix_{(2m+1) \cdot 2^k - 1} + \dots + 2ix_{(m+1) \cdot 2^{k+1} - 2} = 0,$$

d. h.

$$x_0 = 0, \quad x_2 = 0, \quad x_4 = 0, \dots; \quad x_0 + 2x_1 + x_2 = 0, \quad x_4 + 2x_5 + x_6 = 0, \dots;$$

$$x_0 + 2x_1 + 3x_2 + 4x_3 + 3x_4 + 2x_5 + x_6 = 0, \dots, \dots,$$

woraus $x_0 = x_1 = x_2 = \dots = 0$, $f = 0$ folgt.

Es bleibt noch übrig, die $b_{m,k/n,l} = 2^{-\frac{1}{2}(k+l)-1}(\omega_{m,k}, \bar{R}\omega_{n,l})$ zu berechnen, was ohne weiteres geht, da wir $\omega_{m,k}, \bar{R}\omega_{n,l}$ mit den φ_v ausgedrückt haben. Für $m \neq n$, $k = l$ treten keine gemeinsamen φ_v auf, es kommt also 0 heraus, für $m = n$, $k = l$ überzeugt man sich direkt von dem selben. Für $k > l$ und n außerhalb des Intervalles $m \cdot 2^{k-l}, \dots, (m+1) \cdot 2^{k-l} - 1$ fehlen wieder die gemeinsamen φ_v , also 0; liegt n im genannten Intervalle, so ist es $m \cdot 2^{k-l} + \varrho$ oder $= (m+1) \cdot 2^{k-l} - \varrho - 1$ ($\varrho = 0, 1, \dots, 2^{k-l-1} - 1$), und dann berechnet man:

$$(\omega_{m,k}, \bar{R}\omega_{n,l}) = \begin{cases} \frac{(2\varrho+1) \cdot 2^l - 1}{2 \cdot 2^{l+1}} - i(2\alpha + 1) - \frac{(2\varrho+1) \cdot 2^{l+1} - 1}{(2\varrho+1) \cdot 2^l} - i(2\alpha + 1) \\ \quad = -i \cdot 2^l \cdot (-2 \cdot 2^l) = i \cdot 2^{2l+1}, \text{ bzw.} \\ \frac{(2\varrho+1) \cdot 2^l - 1}{(2\varrho+1) \cdot 2^l} - i(2\alpha + 1) - \frac{(2\varrho+1) \cdot 2^l - 1}{2 \cdot 2^{l+1}} - i(2\alpha + 1) \\ \quad = -i \cdot 2^l \cdot 2 \cdot 2^l = -i \cdot 2^{2l+1}. \end{cases}$$

Für $k < l$ muß es (weil $b_{mk/nl}$ Hermitesch ist) nach Vertauschen vom m, n (und k, l) entsprechend $0, -i \cdot 2^{2k+1}, i \cdot 2^{2k+1}$ sein. Und den $0, \pm i \cdot 2^{2k+1}, \pm i \cdot 2^{2l+1}$ entspricht bei $b_{mk/nl}$ $0, \pm i \cdot 2^{\frac{1}{2}(3k-l)}, \pm i \cdot 2^{\frac{1}{2}(3l-k)}$.

Wir führen nun die Indizierung m, k ($= 0, 1, 2, \dots$) auf eine mit p ($= 1, 2, \dots$) zurück, durch $p = (2m + 1) \cdot 2^k$. Dann wird, wie man sich leicht überlegt:

Satz 1.** *Im oben beschriebenen vollst. norm. orth. System $\psi_1, \psi_2, \dots$ gibt es eine H. quadr. Matrix $B = \{b_{pq}\}$, so daß $\bar{R}$ der abg. lin. H. O. von $B, \{\psi_1, \psi_2, \dots\}$ ist. Die b_{pq} ($p, q = 1, 2, \dots$) sind so definiert:*

Wir teilen die Zahlenreihe folgendermaßen ein: Das Intervall $m \cdot 2^{k+1} + 1, \dots, (m+1) \cdot 2^{k+1}$ ist eine Zelle, $m \cdot 2^{k+1} + 1, \dots, (2m+1) \cdot 2^k$ und $(2m+1) \cdot 2^k + 1, \dots, (m+1) \cdot 2^{k+1}$ ihre erste bzw. zweite Hälfte, $(2m+1) \cdot 2^k$ ihre Mitte, k ihr Grad. Jede Zahl $p = 1, 2, \dots$ ist dann Mitte einer einzigen Zelle, ihrer Zelle.

Es ist:

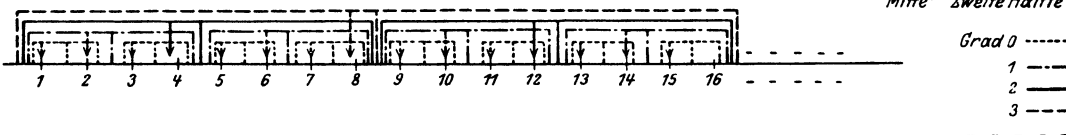
$$b_{pq} \begin{cases} = 0 & \text{wenn von den Zellen von } p, q \text{ keine echter Teil der anderen ist.} \\ = \pm i \cdot 2^{\frac{1}{2}(3k-l)} & \text{wenn das der Fall ist, und } k, l \text{ der Grad der kleineren bzw. größeren Zelle ist.} \end{cases}$$

Das Vorzeichen $\pm$ bestimmt sich so:

	sie liegt in der ersten Hälfte	sie liegt in der zweiten Hälfte
p -s Zelle ist kleiner	—	+
q -s Zelle ist kleiner	+	—

(Man beachte: B ist i -mal eine schiefsymmetrische Matrix.)

Zelle
Erste Hälfte | Mitte | Zweite Hälfte



3. Eine bessere Übersicht gewinnen wir über $\bar{R}$, wenn wir es in Funktionenräumen realisieren. So sei Ω das Intervall $0, 1$, wir bilden seinen Funktionenraum: die Menge aller (nach Lebesgue meßbaren) Funktionen $f(x)$ des Intervalles $0 < x < 1$, mit endlichem $\int_0^1 |f(x)|^2 dx$ —

nach Anhang I ist dies ein Hilbertscher Raum, isomorph $\mathfrak{H}$. Er heiße $\mathfrak{H}^*$.

Die $e^{2n\pi i \cdot x}$, $n = 0, \pm 1, \pm 2, \dots$ bilden bekanntlich ein vollst. norm. orth. System in $\mathfrak{H}^*$. Sei $\mathfrak{F}$ die von den $e^{2n\pi i \cdot x}$, $n = 0, 1, 2, \dots$ aufge-

spannte abg. lin. M., $\mathfrak{S}$ ist unendlichvioldimensional, also ebenfalls ein Hilbertscher Raum (in diesem $\mathfrak{S}$ wollen wir $\bar{R}$ realisieren).

$$Uf(x) = e^{2\pi i \cdot x} f(x)$$

ist ein längentreuer (ja unitärer) Operator in $\mathfrak{S}^*$, der $\mathfrak{S}$ auf ein Teil von sich selbst abbildet. Wenn wir $\varphi_\nu(x) = e^{2\nu\pi i \cdot x}$ ($\nu = 0, 1, 2, \dots$) setzen so sind die $\varphi_0, \varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System in $\mathfrak{S}$, und es ist $U\varphi_\nu = \varphi_{\nu+1}$; d. h. in $\mathfrak{S}$ stimmt U mit $\bar{U}$ überein, und $\bar{R}$ ist seine Cayley-sche Transformierte. Also: $\bar{R}f(x)$ hat Sinn ($f(x)$ aus $\mathfrak{S}$!) wenn $f(x) = \varphi(x) - U\varphi(x) = (1 - e^{2\pi i \cdot x})\varphi(x)$, $\varphi(x)$ aus $\mathfrak{S}$, ist, und zwar ist es dann $= i(\varphi(x) + U\varphi(x)) = i(1 + e^{2\pi i \cdot x})$. Oder auch: $\bar{R}f(x)$ hat Sinn, wenn $f(x), \frac{f(x)}{1 - e^{2\pi i \cdot x}}$ beide zu $\mathfrak{S}$ gehören, und zwar ist es dann $= i \frac{1 + e^{2\pi i \cdot x}}{1 - e^{2\pi i \cdot x}} f(x) = -\operatorname{ctg} \pi x \cdot f(x)$.

$f(x)$ gehöre zu $\mathfrak{S}$. Daß $\frac{f(x)}{1 - e^{2\pi i \cdot x}}$ zu $\mathfrak{S}$ gehört, bedeutet: sein Absolutwertquadratintegral ist endlich, und es ist zu allen $e^{2\pi i \cdot nx}$, $n = -1, -2, \dots$ orth. Hieraus folgt die Endlichkeit von $\int_0^1 \operatorname{ctg}^2 \pi x |f(x)|^2 dx$ (weil ja $\bar{R}f(x) = -\operatorname{ctg} \pi x f(x)$ ist), und wir wollen zeigen, daß diese umgekehrt zur Folge hat, daß $\frac{f(x)}{1 - e^{2\pi i \cdot x}}$ zu $\mathfrak{S}$ gehört ($f(x)$ aus $\mathfrak{S}$!).

Zunächst ist wegen $\left| \frac{1}{1 - e^{2\pi i \cdot x}} \right|^2 = \left| \frac{\operatorname{ctg} \pi x + i}{2i} \right|^2 \leq \frac{1}{2} \operatorname{ctg}^2 \pi x + \frac{1}{2}$ auch $\int_0^1 \left| \frac{f(x)}{1 - e^{2\pi i \cdot x}} \right|^2 dx$ endlich, also alle $\int_0^1 \left| \frac{f(x)}{1 - \varrho e^{2\pi i \cdot x}} \right|^2 dx$ ($0 < \varrho < 1$) absolut gleichmäßig integrierbar. Also ist für $0 < \varrho < 1$, $\varrho \rightarrow 1$

$$\int_0^1 \frac{f(x)}{1 - \varrho e^{2\pi i \cdot x}} e^{-2n\pi i \cdot x} dx \rightarrow \int_0^1 \frac{f(x)}{1 - e^{2\pi i \cdot x}} e^{-2n\pi i \cdot x} dx.$$

Wir haben noch zu zeigen, daß die rechte Seite für alle $n = -1, -2, \dots$ verschwindet, es genügt also dies für die linke (und alle $0 < \varrho < 1$!) zu beweisen. Nun ist $\frac{1}{1 - \varrho e^{2\pi i \cdot x}} = \sum_{\nu=0}^{\infty} \varrho^\nu e^{2\nu\pi i \cdot x}$, und die Reihe konvergiert gleichmäßig in x (ϱ fest!), daher ist

$$\begin{aligned} \sum_{\nu=0}^N \varrho^\nu \int_0^1 f(x) e^{-2(n-\nu)\pi i \cdot x} dx &= \int_0^1 f(x) \left(\sum_{\nu=0}^N \varrho^\nu e^{2\nu\pi i \cdot x} \right) e^{-2n\pi i \cdot x} \\ &\rightarrow \int_0^1 \frac{f(x)}{1 - \varrho e^{2\pi i \cdot x}} e^{-2n\pi i \cdot x} dx \end{aligned}$$

(für $N \rightarrow \infty$). Da $n = -1, 2, \dots$, $\nu = 0, 1, 2, \dots$ also $n - \nu = -1, -2, \dots$

ist und $f(x)$ selbst zu $\mathfrak{S}$ gehört, ist die linke Seite stets 0, also verschwindet auch die rechte.

Nun können wir beweisen:

Satz 2**. $\mathfrak{S}^*$ sei der soeben beschriebene Hilbertsche Funktionenraum, $\mathfrak{S}_k^\pm$ die von allen $e^{2n\pi i \cdot x}$, $n \geq k$ bzw. $\leq k$ ($n = 0, \pm 1, \pm 2, \dots$, ebenso k) aufgespannte abg. lin. M. — die $\mathfrak{S}_k^\pm$ sind unendlichvieldimensional, also wieder Hilbertsche Räume.

R sei der folgende Operator in $\mathfrak{S}^*$: wenn $f(x)$, $-\operatorname{ctg} \pi x \cdot f(x)$ zu $\mathfrak{S}^*$ gehören, (d. h. $\int_0^1 |f(x)|^2 \cdot dx$, $\int_0^1 \operatorname{ctg}^2 \pi x \cdot |f(x)|^2 \cdot dx$ endlich sind), so sei $Rf(x)$ sinnvoll. und zwar $= -\operatorname{ctg} \pi x \cdot f(x)$.

Wenn $f(x)$ in einem der $\mathfrak{S}_k^\pm$ liegt und $Rf(x)$ Sinn hat, so liegt auch $Rf(x)$ in diesem $\mathfrak{S}_k^\pm$; so können wir R in jedem $\mathfrak{S}_k^\pm$ zu einem H. O. dortselbst einschränken⁷⁷⁾. In den $\mathfrak{S}_k^+$ erhalten wir so $\bar{R}$, in den $\mathfrak{S}_k^-$ — $\bar{R}$.

Beweis. Für $\mathfrak{S}_0^+ = \mathfrak{S}$ haben wir dies bereits bewiesen. Die unitäre Abbildung $f(x) \Rightarrow e^{2k\pi i \cdot x} \cdot f(x)$ führt $\mathfrak{S}^*$, R , $\mathfrak{S}_0^+$ in $\mathfrak{S}^*$, R , $\mathfrak{S}_k^+$ über, die unitäre Abbildung $f(x) \Rightarrow f(1-x)$ $\mathfrak{S}^*$, R , $\mathfrak{S}_{-k}^+$ in $\mathfrak{S}^*$, $-R$, $\mathfrak{S}_k^-$ — also gelten die Behauptungen für alle $\mathfrak{S}_k^\pm$.

Man zeigt übrigens leicht, daß dieses R in $\mathfrak{S}^*$ ein hypermax. H. O. ist.

4. Wir können jedem $f(x) = \sum_{n=0}^{\infty} a_n e^{2n\pi i \cdot x}$ ($\sum_{n=0}^{\infty} |a_n|^2$ endlich) von $\mathfrak{S}_0^+$ ⁷⁸⁾, die (wegen $\sum_{n=0}^{\infty} |a_n|^2$ endlich, also $a_n \rightarrow 0$) im Einheitskreise analytische Funktion $F(z) = \sum_{n=0}^{\infty} a_n z^n$ zuordnen. Da

$$\sum_{n=0}^{\infty} |a_n|^2 = \lim_{r \rightarrow 1} \sum_{n=0}^{\infty} r^n |a_n|^2 = \lim_{r \rightarrow 1} \int_{|z|=r} |F(z)|^2 \frac{dz}{z}$$

ist, bilden die im Einheitskreise analytischen Funktionen $F(z)$ mit beschränktem $\int_{|z|=r} |F(z)|^2 \frac{dz}{z}$ ($r \rightarrow 1$) einen Hilbertschen Raum. Wegen

$-\operatorname{ctg} \pi x = i \frac{e^{2\pi i \cdot x} + 1}{e^{2\pi i \cdot x} - 1}$ ist hierin $\bar{R}$ durch $\bar{R} F(z) = i \frac{z+1}{z-1} F(z)$ definiert (falls dieses wieder zur genannten Klasse gehört). — Eine andere be-

⁷⁷⁾ Aber dies ist noch keineswegs Reduzierbarkeit von R durch $\mathfrak{S}_k^\pm$! Denn mit Rf braucht ja nicht auch $RP_{\mathfrak{S}_k^\pm} f$ sinnvoll zu sein.

⁷⁸⁾ Die Reihe $\sum_{n=0}^{\infty} a_n e^{2n\pi i \cdot x}$ konvergiert (im Sinne der von uns stets verwendeten Metrik) „im Mittel“: $\int_0^1 \left| f(x) - \sum_{n=0}^N a_n e^{2n\pi i \cdot x} \right|^2 dx \rightarrow 0$ (für $N \rightarrow \infty$); natürlich nicht notwendig „punktweise“.

merkenswerte Realisation von $\bar{R}$ gewinnen wir, wenn wir die im Intervalle $0, \infty$ definierten Funktionen $f(x)$ (mit endlichem $\int_0^\infty |f(x)|^2 dx$) betrachten: es entspricht dann dem Differentialoperator $i \frac{d}{dx} \dots$ mit der Randbedingung $f(0) = 0$. Wir gehen aber hierauf nicht näher ein ⁷⁹⁾.

Es bleibt noch übrig, die bei Satz 40 benützten Eigenschaften von $\bar{R}$ zu verifizieren. Erstens: $\bar{R}$ ist nicht überall sinnvoll und nimmt nicht jeden Wert an. Wir realisieren es nach Satz 2** in $\mathfrak{S}_0^+$, da $\int_0^1 \operatorname{ctg}^2 \pi dx$, $\int_0^1 \operatorname{tg}^2 \pi x dx$ beide $= \infty$ sind, ist $\bar{R}$ für $f(x) = 1$ sinnlos, und nimmt diesen Wert nie an. Zweitens: $\bar{R}$ ist weder nach oben noch nach unten halbbeschränkt. Wir realisieren es nach Satz 1** und setzen $f = \psi_{2^r} \pm i \psi_{2^r+1}$, dann ist $(f, \bar{R}f) = \mp 2^{r+\frac{1}{2}}$ und $|f|^2 = 2$ — d. h. $(f, \bar{R}f) > C \cdot |f|^2$ bzw. $< C \cdot |f|^2$ bei beliebigem C erreicht, wenn $2^{r-\frac{1}{2}} > |C|$ gewählt wird.

⁷⁹⁾ Sie wird in der in Anm. ²⁸⁾ genannten Arbeit erörtert werden.

Zur Algebra der Funktionaloperationen und Theorie der normalen Operatoren

Einleitung.

1. Die vorliegende Arbeit zerfällt in zwei, im wesentlichen unabhängige, Teile. Der erste (§§ I—III) ist der Untersuchung der linearen und beschränkten Operatoren (d. h. Matrizen) des Hilbertschen Raumes $\mathfrak{H}$ gewidmet, indem die algebraischen Eigenschaften des von ihnen gebildeten (nicht-kommutativen) Ringes $\mathcal{B}$ betrachtet werden. Den Gegenstand des zweiten Teiles hingegen bilden diejenigen, nicht notwendig überall (in $\mathfrak{H}$) sinnvollen und beschränkten Operatoren, die die sogenannte Hilbertsche Spektraldarstellung mit komplexen Eigenwerten zulassen (vgl. die ausführlichere Explizierung dieser Begriffe im § 4 der Einleitung). Dies sind die als „normal“ zu bezeichnenden Operatoren, die bisher nur im Beschränkten betrachtet wurden¹⁾, und für die wir eine neue allgemeinere Definition geben werden (vgl. am vorhin angeführten Orte).

Ehe wir diese Dinge genauer auseinandersetzen, sei an die Definition des (komplexen) Hilbertschen Raumes $\mathfrak{H}$ erinnert. Man kann ihn als die Menge aller Folgen komplexer Zahlen $\{x_1, x_2, \dots\}$ mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$ realisiert denken²⁾; wir bezeichnen seine Elemente mit $f, g, \dots, \varphi, \psi, \dots$.

¹⁾ Bis vor kurzem nur im Vollstetigen, vgl. den Encyklopädie-Artikel von Hellinger und Toeplitz, Encykl. d. math. Wiss. II. C. 13, Seite 1562. Vgl. auch Anm. ^{2b)}.

²⁾ Nach dem bekannten Satze von Fischer und F. Riesz ebensogut auch als Menge aller Funktionen $f(x)$ des Intervalles $a < x < b$ mit endlichem $\int_a^b |f(x)|^2 dx$ ($a < b$, jedes kann auch unendlich sein); oder aller Funktionen $f(P)$, wobei P die Einheitskugel-Oberfläche durchläuft, mit endlichem $\iint |f(P)|^2 d\sigma$ ($\iint \dots d\sigma$ das Integral über die Einheitskugel-Oberfläche), usw. Vgl. z. B. die in Anm. ²⁾ zitierte Arbeit insbesondere Kap. I und Anhang I sowie Einleitung, zum Satz von Fischer und F. Riesz etwa F. Riesz, Göttinger Nachr. 1907, S. 210—273.

Mit diesen kann man wie mit Vektoren rechnen, so daß der Sinn von Bildungen wie $f \pm g$, af ohne weiteres einleuchtet³⁾. Ferner gibt es ein „inneres Produkt“ (f, g) wie bei Vektoren⁴⁾, welches das Definieren des „absoluten Wertes“ im Hilbertschen Raume ermöglicht: $|f| = \sqrt{(f, f)}$, und das Definieren der Entfernung $|f - g|$.⁵⁾ Hierdurch wird der Hilbertsche Raum $\mathfrak{H}$ zu einem topologischen (ja metrischen) linearen Raume⁶⁾, in welchem man geometrisch-topologische Redeweisen, wie lineare Mannigfaltigkeit, Häufungspunkt, Stetigkeit usw., ohne weiteres verwenden darf. Zur konsequenten Kennzeichnung von $\mathfrak{H}$ auf dieser Grundlage vergleiche man etwa mit einer früheren Arbeit des Verf.⁷⁾, deren Bezeichnungen und Begriffsbildungen hier auch sonst verwendet werden sollen. Vgl. hierzu besonders die in Anm. 2) genannten Teile derselben; immerhin werden wir die zur Anwendung kommenden Abschnitte von E. stets genau zitieren.

Die für unsere Zwecke wichtigsten Definitionen aus E. sind (Abkürzungen nach Anm. 7)): Ein O. ist eine Funktion, deren Wertevorrat und Definitionsbereich Teilmengen von $\mathfrak{H}$ sind — wir bezeichnen die O. mit $A, B, \dots, R, S, \dots$, den Wert von A an der Stelle f mit Af . (Af braucht also keineswegs überall in $\mathfrak{H}$ sinnvoll zu sein.) Er ist lin., wenn mit Af, Ag auch $A(af)$, $A(f + g)$ sinnvoll sind (d. h. der Definitionsbereich

3) Es wird definiert: $\{x_1, x_2, \dots\} \pm \{y_1, y_2, \dots\} = \{x_1 \pm y_1, x_2 \pm y_2, \dots\}$, $a\{x_1, x_2, \dots\} = \{ax_1, ax_2, \dots\}$, wenn $\mathfrak{H}$ als „Folgenraum“ der $\{x_1, x_2, \dots\}$ (mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$) interpretiert ist. Im „Funktionsraume“ der $f(x)$ ($\int_a^b |f(x)|^2 dx$ endlich) ist entsprechend $(f \pm g)(x) = f(x) \pm g(x)$, $(af)(x) = af(x)$; analog für andere Beispiele. Vgl. das Zitat von Anm. 2).

4) Es wird definiert: $(\{x_1, x_2, \dots\}, \{y_1, y_2, \dots\}) = \sum_{n=1}^{\infty} x_n \bar{y}_n$ (die Reihe konvergiert absolut) bzw. $(f(x), g(x)) = \int_a^b f(x) \bar{g}(x) dx$, usw. Vgl. wie oben.

5) Somit ist $|\{x_1, x_2, \dots\}| = \sqrt{\sum_{n=1}^{\infty} |x_n|^2}$ bzw. $|f(x)| = \sqrt{\int_a^b |f(x)|^2 dx}$ (leider ist hier die Bezeichnung mißverständlich, auf der linken Seite steht der Absolutwert der Funktion $f(x)$, rechts diejenigen ihrer Werte — bei unseren Überlegungen in $\mathfrak{H}$ wird aber so etwas nie vorkommen), usw. Vgl. wie oben.

6) Vgl. Hausdorff, Umgebungsaxiome und topologisch-lineare Räume: Grundzüge der Mengenlehre, I. Auflage, Leipzig und Berlin 1914.

7) „Zur allgemeinen Eigenwerttheorie Hermitescher Funktionaloperatoren“, Math. Annalen 102, 1 (1929), S. 49–131. Diese im folgenden mehrfach zu zitierende Arbeit soll stets mit E. bezeichnet werden. Von Wichtigkeit sind einige Abkürzungen aus E., die wir verwenden werden; sie seien hier zusammengestellt: Abgeschlossen = abg., beschränkt = beschr., Fortsetzung = Forts., Hermitesch = H., hypermaximal = hypermax., linear = lin., lineare Mannigfaltigkeit = lin. M., maximal = max., normiert = norm., Operator = O., orthogonal = orth., Projektion = P., vollständig = vollst., Zerlegung der Einheit = Z. d. E. Ferner werden wir den Begriff „vertauschbar“ (vgl. Anm. 16)) mit vert. abkürzen.

eine lin. M. ist⁹⁾), und zwar gleich aAf bzw. $Af + Ag$. Er heißt *abg.*, wenn er die folgende Eigenschaft besitzt: Wenn alle Af_n ($n = 1, 2, \dots$) sinnvoll sind, und $f_n \rightarrow f$, $Af_n \rightarrow f^*$ (für $n \rightarrow \infty$) gilt, so ist Af sinnvoll und gleich f^* ⁹⁾).

Stetig ist ein O., wenn er es im Sinne der üblichen Definition in seinem ganzen Definitionsbereiche ist: Wenn Af_0 Sinn hat, so existiert zu jedem $\varepsilon > 0$ ein $\delta > 0$, so daß aus $|f - f_0| < \delta$ $|Af - Af_0| < \varepsilon$ folgt (wenn Af Sinn hat). Man sieht sofort: Für einen *abg. stetigen*[†]O. ist der Definitionsbereich notwendig eine *abg. Menge* — ist der O. auch lin., so ist der Definitionsbereich also eine *abg. lin. M.*¹⁰⁾.

In der ersten Hälfte dieser Arbeit werden wir uns, wie schon angekündigt, mit überall sinnvollen, lin., stetigen O. beschäftigen, die (in Anlehnung an Hilbert) beschr. heißen sollen. (Vgl., auch für das Folgende, den Anhang I von E.) Für diese überall sinnvollen O. sind die einfachsten Rechenspezies so zu definieren: Wenn A, B beschr. sind, so sind es aA , $A \pm B$, AB auch, und zwar gilt

$$(aA)f = a \cdot Af, \quad (A \pm B)f = Af \pm Bf, \quad (AB)f = A(Bf).$$

Es gelten die Regeln der Algebra, mit zwei Ausnahmen: Es gibt Nullteiler, und die Multiplikation ist nicht kommutativ¹¹⁾. Die Menge $\mathcal{B}$ aller beschr. O. ist also ein Ring — mit Nullteilern und nicht kommutativ. Eine weitere wichtige Operation ist $*$, die so definiert ist: Zu jedem beschr. A gibt es genau ein beschr. A^* mit

$$(Af, g) = (f, A^*g), \quad (f, Ag) = (A^*f, g).^{12)}$$

⁹⁾ D. h. mit f, g auch af , $f + g$ enthält. Eine *abg. lin. M.* muß auch noch *abg.* sein, d. h. alle ihre Häufungspunkte enthalten.

⁹⁾ Diese stetigkeitsähnliche Eigenschaft ist in kompakten Räumen der Stetigkeit gleichwertig, aber in § viel schwächer; vgl. E. Anm. ²⁹⁾.

¹⁰⁾ Für lin. O. kann die Stetigkeit auf mehrere, gleichwertige Formulierungen zurückgeführt werden, vgl. E. Satz 12. Eine von diesen ist: Es gebe ein festes c , so daß aus $|f| \leq 1$ $|Af| \leq c$ folgt.

¹¹⁾ 0, 1 sind durch $0f = 0$, $1f = f$ zu definieren.

¹²⁾ Jede der beiden Relationen folgt, durch Vertauschen von f, g und Nehmen der Komplex-Konjugierten, aus der anderen. — Der Operation $*$ entspricht bei Matrizen das Nehmen der Komplex-Konjugiert-Transponierten. D. h. wenn § als „Folgenraum“ der $\{x_1, x_2, \dots\}$ $\left(\sum_{n=1}^{\infty} |x_n|^2 \text{ endlich}\right)$ realisiert wird und A die Matrix $\{a_{\mu\nu}\}$ hat (also wenn stets

$$A\{x_1, x_2, \dots\} = \{y_1, y_2, \dots\},$$

$$y_\nu = \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu$$

gilt — analog wie im Anhang II von E.), so hat A^* die Matrix $\{\overline{a_{\nu\mu}}\}$. Vgl. E. Anhang II.

[†] This line should read: Für einem *abg. gleichmassigen stetigen*

Man verifiziert leicht die Regeln:

$$0^* = 0, \quad 1^* = 1, \quad (aA)^* = \bar{a}A^*, \quad (A \pm B)^* = A^* \pm B^*, \quad (AB)^* = B^*A^*$$

Bei den unbeschr. (genauer: nicht notwendig beschr.) O., die den Gegenstand der zweiten Hälfte dieser Arbeit bilden, müssen wir etwas anders vorgehen. Für solche O. R, S ist zunächst bei der Definition von $aR, R \pm S, RS$ zu beachten, daß die Definitionsbereiche dieser O. dadurch eingeengt werden, daß Rf, Sf nicht überall sinnvoll sein müssen. Damit $(R + S)f$ Sinn habe, müssen ja Rf, Sf sinnvoll sein, ebenso bei $(R - S)f$ — und bei RSf muß Sf und $R(Sf)$ Sinn haben. Bei R^* ist es sogar fraglich, ob ein solcher O. existiert. Für uns ist aber gerade die Operation $*$ von großer Wichtigkeit, wir werden darum gegebenenfalls gleich vom fertigen Paare R, R^* ausgehen, indem wir festsetzen: Ein konjugiertes O.-Paar (kurz: konj. O.-Paar) wird von zwei O. R, R^* mit demselben Definitionsbereich gebildet, für die (soweit beide Seiten sinnvoll sind) stets

$$(Rf, g) = (f, R^*g), \quad (f, Rg) = (R^*f, g)$$

gilt (vgl. Anm. ¹²⁾). Außerdem verlangen wir, analog wie bei den H. O. und aus denselben Gründen (vgl. E. Einleitung V sowie Definition 6 und Anm. ³⁰⁾), daß der Definitionsbereich die abg. lin. M. $\S$ aufspanne. Diese Definition wird sich als die richtige Verallgemeinerung der Operation $*$ in $\mathcal{B}$ erweisen, die H. O. sind in ihr als Spezialfall $R = R^*$ mit erfaßt.

2. Zunächst sei einiges über die Gesichtspunkte der Untersuchung von $\mathcal{B}$ gesagt. Wir fassen diesen Ring als ein hyperkomplexes Zahlensystem auf: In der Tat ist es wohl das einfachste derartige System mit unendlich vielen Einheiten — d. h. das einfachste, welches über die im allgemeinen untersuchten und im wesentlichen durch die Gültigkeit der sogenannten „Teilerkettensätze“ gekennzeichneten ¹³⁾ Systeme hinausgeht. Während aber in den hyperkomplexen Systemen mit endlich vielen Einheiten (bzw. den dies ersetzenden „Teilerkettensätzen“) die Topologie so trivial ist, daß sie gar keine nennenswerte Rolle hat, spielt in unserem unendlichen System $\mathcal{B}$ die Topologie eine entscheidende Rolle. So ist es z. B. bei der definitorischen Festlegung dessen, welche Teilmengen $\mathcal{M}$ von $\mathcal{B}$ auch Ringe (in der Terminologie der Anm. ¹³⁾ zitierten Arbeiten: „Ordnungen“) sind — wenn wir nur verlangten, daß $\mathcal{M}$ mit A, B auch $aA, A \pm B, AB$ enthält, würden wir uns in ein unentwirrbares Dickicht pathologischer Bildungen verlieren —, es muß nämlich die Forderung gestellt werden, daß $\mathcal{M}$ im Sinne einer (noch zu formulierenden) Topologie von $\mathcal{B}$ abg. sei.

¹³⁾ Vgl. E. Noether, *Math. Annalen* **96**, 1 (1926), S. 26–61, insbesondere § 2; ferner E. Artin, *Abh. d. Math. Sem. d. Hamb. Univ.* **5**, 3 (1927), S. 251–260.

Übrigens werden wir noch eine vereinfachende Forderung an die Ringe $\mathcal{M}$ stellen: mit A soll auch A^* dazugehören. Dies ist eine bedeutsame Einschränkung, die es ermöglicht, den Komplikationen der Elementarteilertheorie aus dem Wege zu gehen¹⁴⁾. Es ist wohl motiviert, diese Schwierigkeiten zunächst auszuschalten; die Theorie von $\mathcal{B}$ liefert, gegenüber dem endlichvioldimensionalen Falle¹⁵⁾, ohnehin genug neue Komplikationsmöglichkeiten.

Aus dem oben Gesagten geht hervor, daß wir uns zunächst um die Topologie von $\mathcal{B}$ kümmern müssen. Dabei wirkt es störend, daß es mehrere Möglichkeiten des Topologisierung (d. h. des Definierens des Umgebungsbegriffes) in $\mathcal{B}$ gibt, die für uns in Betracht kommen. Diese werden wir im § I genau diskutieren; es genüge hier, wenn wir sagen, daß bei Ringen $\mathcal{M}$ in $\mathcal{B}$ die Abg. der sogenannten schwachen Topologie verlangt werden soll.

Da der Durchschnitt beliebig vieler Ringe wieder ein Ring ist, ist insbesondere der Durchschnitt aller $\mathcal{M}$ umfassenden Ringe einer ($\mathcal{M}$ sei irgendeine Teilmenge von $\mathcal{B}$), er ist offenbar als kleinster $\mathcal{M}$ enthaltender Ring eindeutig gekennzeichnet — wir nennen ihn $R(\mathcal{M})$. Eine weitere wichtige Begriffsbildung ist die folgende: Sei $\mathcal{M}$ eine Teilmenge von $\mathcal{B}$; die Menge aller A aus $\mathcal{B}$, für die A und A^* mit jedem B von $\mathcal{M}$ vert. sind¹⁶⁾, heiße $\mathcal{M}'$. Diese Operation ' kann iteriert werden, dadurch entstehen aus $\mathcal{M}$ der Reihe nach die Mengen $\mathcal{M}', \mathcal{M}'', \mathcal{M}''', \dots$. Sie hat eine Reihe einfacher Eigenschaften, so werden wir u. a. zeigen: Es ist stets $\mathcal{M} < \mathcal{M}''$,¹⁷⁾ aus $\mathcal{M} < \mathcal{N}$ folgt $\mathcal{M}' > \mathcal{N}'$, es gilt allgemein $\mathcal{M}' = \mathcal{M}''' = \mathcal{M}^V = \dots$, $\mathcal{M}'' = \mathcal{M}^{IV} = \mathcal{M}^{VI} = \dots$.

Die beiden Operationen $R(\mathcal{M})$, $\mathcal{M}'$ ermöglichen zusammen eine einfache Kennzeichnung der algebraischen Verhältnisse in $\mathcal{B}$.

3. Ein O. A heißt normal, wenn A, A^* vert. sind (vgl. dazu E. Anhang I). Ein Ring $\mathcal{M}$ heißt Abelsch, wenn alle seine Elemente miteinander vert. sind. Man zeigt leicht: Ein Ring $\mathcal{M}$ ist dann und nur dann Abelsch, wenn alle seine Elemente normal sind, und der Ring $R(A)$ ist es dann und nur dann, wenn A normal ist. Wir werden nun die Umkehrung beweisen: Jeder Abelsche Ring $\mathcal{M}$ kann durch einen normalen O. (sogar durch einen H. O.) erzeugt werden — d. h. es gibt ein solches A mit $\mathcal{M} = R(A)$.

¹⁴⁾ Für Matrizen in endlichvioldimensionalen Räumen hat E. Fischer diese Bedingung als erster mit Erfolg eingeführt.

¹⁵⁾ Jedes hyperkomplexe System mit endlich vielen Einheiten kann ja durch Matrizen (also Operatoren) eines endlichvioldimensionalen Raumes dargestellt werden.

¹⁶⁾ Zwei A, B aus $\mathcal{B}$ sind vert., wenn $AB = BA$ ist.

¹⁷⁾ $\mathcal{M} < \mathcal{N}$ oder $\mathcal{N} > \mathcal{M}$ bedeutet: $\mathcal{M}$ ist Teilmenge von $\mathcal{N}$.

Eine weitere, bei Abelschen Ringen eintretende Vereinfachung ist die folgende: Sei $\mathcal{M}$ eine Teilmenge von $\mathcal{B}$, $\mathbf{R}(\mathcal{M})$ der Ring von $\mathcal{M}$ — man zeigt leicht, daß $\mathbf{R}(\mathcal{M})$ entsteht, indem man zuerst aus den Elementen von $\mathcal{M}$ alle durch (endlichvielmaliges) Anwenden der Operationen aA , $A+B$, AB und A^* entstehenden O. und deren Menge $\mathbf{r}(\mathcal{M})$ bildet, und dann zu $\mathbf{r}(\mathcal{M})$ alle seine Häufungspunkte hinzunimmt. Dabei ist Häufungspunkt, wie in 2. erwähnt, im Sinne der sogenannten schwachen Topologie zu verstehen. Wenn nun $\mathbf{R}(\mathcal{M})$ Abelsch ist (dies ist dann und nur dann der Fall, wenn alle Elemente von $\mathcal{M}$ miteinander und mit ihren $*$ vert. sind), so werden wir zeigen, daß es genügt eine viel geringere Menge von Häufungspunkten — die sogenannten Limites stark doppelkonvergenter Folgen, vgl. die Diskussion dieser Begriffe im § I — zu $\mathbf{r}(\mathcal{M})$ hinzuzufügen: schon dadurch entsteht ganz $\mathbf{R}(\mathcal{M})$. Dieses Resultat ist insbesondere für die Theorie der Funktionen von normalen (also insbesondere H. und unitären) vert. O. wichtig, allerdings führt dieselbe aus dem Rahmen der vorliegenden Arbeit hinaus. —

Für beliebige (nicht notwendig Abelsche!) Ringe werden wir eine Beziehung zwischen $\mathbf{R}(\mathcal{M})$ und $\mathcal{M}''$ ($\mathcal{M}$ eine beliebige Teilmenge von $\mathcal{B}$) aufstellen. Wir werden insbesondere zeigen (unser weitestgehendes Resultat lautet noch etwas allgemeiner, vgl. § II, insbesondere Satz 5): Wenn 1 zu $\mathcal{M}$ gehört, so ist $\mathbf{R}(\mathcal{M}) = \mathcal{M}''$. Um die Bedeutung dieses Satzes für das als hyperkomplexes Zahlensystem aufzufassende $\mathcal{B}$ einzusehen, vergegenwärtige man sich, was er besagt, wenn der Hilbertsche Raum $\mathfrak{H}$ durch einen endlichviel-, etwa k -dimensionalen Euklidischen Raum ersetzt wird. $\mathcal{M}$ sei speziell eine (endliche oder auch kontinuierliche) Gruppe unitärer Matrizen, dann erkennt man sofort, daß $\mathbf{R}(\mathcal{M})$ die Menge aller Linearaggregate von Matrizen aus $\mathcal{M}$ ist (unter denen es ja höchstens k^2 lin. unabhängige gibt!).

Sei zunächst $\mathcal{M}$ (als Matrizenmenge) irreduzibel, dies bedeutet bekanntlich, daß nur die Matrizen $a1$ mit allen Elementen von $\mathcal{M}$ vert. sind¹⁸⁾, d. h. daß $\mathcal{M}'$ aus ihnen allein besteht. $\mathcal{M}''$ umfaßt also alle Matrizen überhaupt, nach unserem Satze muß dies auch $\mathbf{R}(\mathcal{M})$ tun — somit ist jede Matrix ein Linearaggregat der Matrizen aus $\mathcal{M}$, oder was dasselbe ist (da es k^2 lin. unabhängige Matrizen gibt): es gibt in $\mathcal{M}$ k^2 lin. unabhängige Matrizen. Oder auch: Keine feste lin. Relation zwischen den k^2 Matrizenelementen kann für alle Elemente von $\mathcal{M}$ bestehen. Das ist aber der Burnsidesche Satz über irreduzible Matrizengruppen¹⁹⁾. Wenn wir $\mathcal{M}$ nicht

¹⁸⁾ Vgl. I. Schur, Berl. Ber. 1905, S. 406. Die dort stets vorausgesetzte Endlichkeit von $\mathcal{M}$ ist in diesem Falle bekanntlich unwesentlich.

¹⁹⁾ Vgl. Burnside, Proc. London Math. Soc. (2) 3 (1905), S. 430. Allerdings nur für unitäre Matrizen!

als irreduzibel voraussetzen, erkennt man ebenso mühelos, daß unser Satz der Frobenius - I. Schurschen Verschärfung des Burnsidischen Satzes entspricht²⁰⁾).

Wir haben somit das Analogon desjenigen Satzes bewiesen, der in endlichvielen Dimensionen zur Begründung der unitären Darstellungstheorie von Gruppen dienen kann. Wieweit ein Äquivalent dazu im Hilbertschen Raume auf unseren Satz gegründet werden kann, soll hier nicht untersucht werden.

Mit diesen Ausführungen ist die erste Hälfte der vorliegenden Arbeit erschöpft.

4. Der Gegenstand der zweiten Hälfte (§§ IV—V) sind die am Ende von 1. erwähnten konj. O.-Paare — unabhängig von jeder Beschr.-Annahme. Zunächst gilt es, die normalen O. unter ihnen zu kennzeichnen. Hier entsteht nun eine Schwierigkeit: das Paar R, R^* sollte normal heißen, wenn R und R^* vert. sind — da aber R, R^* nicht überall sinnvoll sein müssen, ist die Definition der Produkte RR^*, R^*R bzw. die Zurückführung der Vert. auf dieselben zunächst unsicher²¹⁾. Daß auch die naheliegende Matrizen-Definition der Vert. versagt, werden wir im Anhang III sehen.

Wir helfen uns wie folgt. Wenn von zwei O. R, A wenigstens A überall sinnvoll ist, so definieren wir die Vert. so: AR sei Forts. von RA — d. h. wenn Rf Sinn hat, so ist auch $R(Af)$ sinnvoll, und zwar gleich $A(Rf)$ ($Af, A(Rf)$ sind ohnehin sinnvoll²²⁾). Daher können wir auch für eine Menge $\mathcal{M}$ ganz beliebiger O. (die nicht einmal lin. und ebenso wenig überall sinnvoll sein müssen) $\mathcal{M}'$ bilden: Es ist die Menge aller A von $\mathcal{B}$, für die A wie A^* mit jedem R von $\mathcal{M}$ vert. ist — $\mathcal{M}'$ ist also immer $\subset \mathcal{B}$. Insbesondere haben wir für jeden O. R die Mengen $(R)', (R)'', \dots$ ²³⁾ (alle $\subset \mathcal{B}!$).

²⁰⁾ Vgl. Frobenius und I. Schur, Berl. Ber. 1906, S. 209.

²¹⁾ Keinesfalls ist für die Vert. von R, S $RS = RS$ kennzeichnend. Sei nämlich Rf nicht überall sinnvoll, dann gilt von $0Rf$ dasselbe, während $R0f$ überall sinnvoll ist — also ist $0R \neq R0$, d. h. R mit 0 nicht vert., was bei einem vernünftigen Vert.-Begriff nicht eintreten darf! — Das Bilden von $AB, A \pm B$ bei nicht überall sinnvollen A, B ist überhaupt eine fragwürdige Angelegenheit. So gibt es z. B. hypermax. H. O. A, B , für die niemals (außer für $f=0$) Af, Bf gleichzeitig sinnvoll sind, also $A+B$ gar nicht zu bilden ist. (Vgl. die demnächst im Journal f. Math. erscheinende Arbeit des Verf. „Zur Theorie der unbeschränkten Matrizen“, wo sehr allgemeine Klassen von Gegenbeispielen angegeben werden.)

²²⁾ Gerade das Beispiel $R, 0$ in Anm. ²¹⁾ legt diese Definition nahe. Wenn ferner E ein P. O. ist, so besagt die Vert. von R, E , daß R durch die zu E gehörige abg. lin. M. $\mathcal{M}$ reduziert wird (vgl. E., Definition 13), wie es vernünftigerweise zu erwarten ist.

²³⁾ Es ist wohl nicht mißverständlich, daß wir, wie schon bei $R(A)$, die Menge mit dem einzigen Element R auch R nennen.

Nun zeigt man leicht (vgl. § II, 1.): Ein A von $\mathcal{B}$ ist dann und nur dann normal, wenn $(A)''$ Abelsch ist. Diese Definition dehnen wir auf beliebige konj. O.-Paare aus: sie sollen normal heißen, wenn $(R)''$ Abelsch ist. Von der Anwendbarkeit dieser Definition auf weite Klassen von O. werden wir uns im Anhang II überzeugen: so ist die im folgenden zu entwickelnde Theorie ohne weiteres auf die (komplexen, unbeschränkten) Laurentschen Matrizen und ihre Verallgemeinerungen²⁴⁾ anwendbar.

Von diesen normalen O. werden wir nun zeigen, daß sie und nur sie eine (und zwar nur eine) Hilbertsche Spektralform mit komplexen Eigenwerten besitzen²⁵⁾. Diese Untersuchungen sind daher wesentlich anders gerichtet als die vorhergehenden Paragraphen und eher als eine Fortsetzung von E. anzusehen.

Wie man sieht, liegen bei normalen O. die Verhältnisse etwas günstiger wie bei H. O.: bei diesen war ja die Hilbertsche Spektralform nicht immer erreichbar²⁶⁾ — und dies ist um so sonderbarer, als im Endlichvioldimensionalen und auch noch im Beschränkten H. ein Spezialfall normal war ($AA^* = A^*A$ folgt aus $A = A^*$)! Dies klärt sich dadurch auf, daß wir zeigen werden: Ein H. O. ist im wesentlichen dann und nur dann normal, wenn er hypermax. ist, und gerade diese H. O. haben eine Hilbertsche Spektralform (vgl. E. Satz 36). Daß in den oben genannten Fällen H. als Spezialfall von normal erscheint, rührt also einfach daher, daß dort alles hypermax. ist.

5. Die Einteilung dieser Arbeit ist die folgende: Im § I werden die verschiedenen Topologisierungsmöglichkeiten von $\mathcal{B}$ angegeben und diskutiert. Im § II führen wir die Operationen $\mathcal{M}'$ und $R(\mathcal{M})$ (Ringbildung) ein und beweisen verschiedene Eigenschaften derselben (u. a. das schon angekündigte Analogon des Burnside-Frobenius-Schurschen Satzes).

²⁴⁾ Vgl. Toeplitz, Math. Annalen 70 (1911), S. 351–376.

²⁵⁾ Zur Präzisierung dieses, die gewöhnliche (reelle) Hilbertsche Spektralform verallgemeinernden Begriffes vgl. E. Anhang II. — Der Begriff der Normalität bei endlichvioldimensionalen Matrizen stammt von Frobenius [Journ. f. Math. 84 (1877), S. 51–54]; er bewies, daß diese und nur diese Matrizen unitär auf die Diagonalforn (mit komplexen Diagonalelementen!) transformiert werden können. Im Encyklopädie-Artikel von Hellinger und Toeplitz (II. C. 13, S. 1562–1563) wurde dasselbe für die vollstetigen O. des Hilbertschen Raumes gezeigt.

In E. Anhang I bewies der Verf., daß alle beschr. normalen O. auf die komplexe Hilbertsche Spektralform gebracht werden können (das ist eben das Äquivalent der Diagonalforn). Hier soll die Theorie vervollständigt und auf unbeschr. O. ausgedehnt werden. — Unabhängig vom Verf., und unter Verwendung einer anderen Methode, hat seither A. Wintner die unitären O. (die ein Spezialfall der beschr. normalen sind) ebenfalls auf die Spektralform gebracht. Vgl. Math. Zeitschr. 30, 1/2 (1929), S. 228–282.

²⁶⁾ Vgl. E., Einleitung VII.

Im § III diskutieren wir die Abelschen Ringe (vgl. das früher darüber Gesagte). Die §§ IV, V sind vom früheren (bis auf die Definitionen) fast unabhängig, ihr Inhalt ist die allgemeine Theorie der normalen Operatoren (ohne Voraussetzung der Realität oder der Beschränktheit) und deren Spektraldarstellung.

Anhang I behandelt die beschr. H. O., Anhang II gibt eine Anwendung der Theorie der normalen Operatoren auf Laurentsche und allgemeinere Matrizen. In Anhang III wird die Unbrauchbarkeit der üblichen Definition der Matrizenvertauschbarkeit im Unbeschränkten gezeigt.

I. Topologisierungen von $\mathfrak{S}$ und $\mathcal{B}$.

1. Ehe wir die Topologie von $\mathcal{B}$ betrachten, wollen wir die etwas einfacheren topologischen Verhältnisse in $\mathfrak{S}$ beschreiben — weniger der späteren Anwendungen halber, als um die Situation in $\mathcal{B}$ exemplifizieren zu können.

Wir hatten bisher (wie in E.) nur die folgende Topologie in $\mathfrak{S}$ betrachtet, die als „starke“ bezeichnet wird²⁷⁾, und die wir nach Hausdorff durch Angabe aller Umgebungen eines willkürlichen Punktes f_0 von $\mathfrak{S}$ kennzeichnen wollen:

Starke Topologie in $\mathfrak{S}$. Sei $\mathcal{U}_1(f_0; \varepsilon)$ die Menge aller f mit $|f - f_0| < \varepsilon$; alle $\mathcal{U}_1(f_0, \varepsilon)$ mit $\varepsilon > 0$ sind die Umgebungen von f_0 .

Diese Umgebungen sind somit die Kugellinneren um f_0 . Daraus, daß $|f - f_0|$ eine Entfernung ist (vgl. E., Definition 4 und Anm. ²⁷⁾) folgt sofort, daß die vier Hausdorffschen Umgebungsaxiome erfüllt sind (die man in jeder Topologie postuliert), und zwar²⁸⁾:

- α) Ein Punkt ist in jeder seiner Umgebungen enthalten.
- β) Der Durchschnitt von zwei Umgebungen eines Punktes umfaßt gewiß noch eine Umgebung desselben.
- γ) Jeder Punkt in einer Umgebung besitzt selbst eine Umgebung, die Teilmenge der ersteren ist.
- δ) Zwei verschiedene Punkte haben immer zwei elementfremde Umgebungen.

²⁷⁾ Vgl. Weyl, Dissertation Göttingen 1908, S. 8—9.

²⁸⁾ Vgl. Hausdorff, Grundzüge der Mengenlehre, I. Auflage, Leipzig und Berlin 1914; daselbst wird auf diesen Umgebungsbegriff die ganze Topologie aufgebaut. Wir erwähnen, wie Häufungspunkte und Limes zu definieren sind. Ein Punkt ist Häufungspunkt einer Menge, wenn jede seiner Umgebungen Punkte der Menge enthält; er ist Limes einer Folge, wenn jede seiner Umgebungen alle Punkte derselben, mit endlich vielen Ausnahmen, enthält. Man erkennt mühelos: Limes sein, heißt Häufungspunkt von jeder unendlichen Teilmenge der Folge sein.

Wir wissen ferner schon aus E., daß die Operationen af , $f \pm g$, (f, g) , $|f|$ in diesem Sinne stetig sind (und zwar $f \pm g$, (f, g) in beiden Variablen gleichzeitig!).

In der starken Topologie gilt das sogenannte erste Hausdorffsche Abzählbarkeitsaxiom²⁹⁾: zu jedem f_0 gibt es eine absteigende Umgebungsfolge derart, daß jede Umgebung von f_0 eine solche aus dieser Folge umfaßt — man nehme die $\mathcal{U}_1(f_0; \frac{1}{n})$, $n = 1, 2, \dots$. Hieraus folgert man mühelos: Zu jedem Häufungspunkte einer Menge (in $\S$) gibt es eine Teilfolge dieser Menge, von der er der Limes ist.

Nun ist es auch üblich, eine andere Topologie in $\S$ einzuführen, die sogenannte „schwache“³⁰⁾. D. h. gewöhnlich definiert man nur die schwachen Limes und nicht die Umgebungen, und zwar so: $f_1, f_2, \dots$ konvergiert schwach gegen f , wenn für jedes feste φ $(f_n, \varphi) \rightarrow (f, \varphi)$. Es ist aber leicht, dies auf einen Umgebungsbegriff zurückzuführen:

Schwache Topologie in $\S$. Sei $\mathcal{U}_2(f_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ die Menge aller f mit $|(f - f_0, \varphi_1)| < \varepsilon, \dots, |(f - f_0, \varphi_s)| < \varepsilon$; alle $\mathcal{U}_2(f_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ mit beliebigen $\varphi_1, \dots, \varphi_s$ (aus $\S$, s beliebig) und $\varepsilon > 0$ sind die Umgebungen von f_0 .

Wieder erkennt man leicht das Erfülltsein von Hausdorffs Axiomen $\alpha)$ bis $\delta)$, sowie die Tatsache, daß ein Limes im Sinne dieser Topologie mit dem oben angegebenen schwachen Limes identisch ist. Die Stetigkeit von af , $f \pm g$ (die letztere in beiden Variablen zugleich) ist evident, die von (f, g) in f — bei festem g — definitorisch, wegen $(f, g) = \overline{(g, f)}$ ist es also auch in g — bei festem f — so. Aber in beiden Variablen zugleich ist es nicht der Fall, denn dann wäre auch $|f| = \sqrt{(f, f)}$ stetig, was nicht zutrifft³¹⁾.

Da $\mathcal{U}_1(f_0; \varepsilon)$ Teilmenge von $\mathcal{U}_2(f_0; \varphi_1, \dots, \varphi_s, \eta)$ ist (z. B. für $\varepsilon = \eta : \text{Max}(|\varphi_1|, \dots, |\varphi_s|)$) ist jeder Häufungspunkt bzw. Limes im starken Sinne auch einer im schwachen — die Umkehrung kann schon wegen des verschiedenen Stetigkeitscharakters von $|f|$ nicht gelten.

Nun gilt in der schwachen Topologie von $\S$ das erste Hausdorffsche Abzählbarkeitsaxiom nicht, denn seine wichtigste Folgerung ist nicht erfüllt: es gibt eine Menge und einen Häufungspunkt derselben, der für keine Teilfolge von ihr Limes ist — d. h. die schwache Abg. kann, im Gegensatz

²⁹⁾ Vgl. am in Anm. ²⁸⁾ a. O., wegen der Separabilität von $\S$ (E., § I, Axiom C) gilt auch das zweite, vgl. ebendort.

³⁰⁾ Weyl (vgl. Anm. ²⁷⁾) definiert sie etwas anders als wir. Vgl. S. 9 a. a. O.

³¹⁾ Denn in jedem $\mathcal{U}_2(f_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ liegen noch f mit $|f| = 1$: es genügt f zu allen $\varphi_1, \dots, \varphi_s$ orth. und mit dem Betrage 1 zu wählen!

zu dem, was man sonst gewohnt ist, nicht auf Konvergenz-Eigenschaften zurückgeführt werden.

Ehe wir das Gegenbeispiel angeben, erwähnen wir noch: wenn $f_1, f_2, \dots$ schwach gegen f konvergiert, so sind die $|f_1|, |f_2|, \dots$, nach einer Bemerkung von E. Schmidt, beschränkt³²⁾. Das Beispiel ist die Menge $\mathfrak{A}$ (wir denken uns $\mathfrak{H}$ als Folgenraum der $\{x_1, x_2, \dots\}$ realisiert) aller $\{x_1, x_2, \dots\}$ für die nur zwei $x_m, x_n \neq 0$ sind, und zwar $x_m = 1, x_n = m$ (sonst m, n beliebig). 0 ist schwacher Häufungspunkt von $\mathfrak{A}$, d. h. bei gegebenen $\{u_1^1, u_2^1, \dots\}, \dots, \{u_1^s, u_2^s, \dots\}$, ε können wir $|\sum_{p=1}^{\infty} u_p^1 x_p| < \varepsilon, \dots, |\sum_{p=1}^{\infty} u_p^s x_p| < \varepsilon$ ($\{x_1, x_2, \dots\}$ aus $\mathfrak{A}$) erzwingen: wir wählen m so, daß $|u_m^1| < \frac{\varepsilon}{2}, \dots, |u_m^s| < \frac{\varepsilon}{2}$ gilt (es ist ja $u_p^1 \rightarrow 0, \dots, u_p^s \rightarrow 0$), und dann n so, daß $|u_n^1| < \frac{\varepsilon}{2m}, \dots, |u_n^s| < \frac{\varepsilon}{2m}$ gilt. Dagegen konvergiert keine Teilfolge von $\mathfrak{A}$ schwach gegen 0: in einer solchen wären die Absolutwerte, d. i. $\sqrt{1+m^2}$, beschränkt, d. h. es kämen nur endlich viele m darin vor — für ein m_0 müßte also

³²⁾ Dies folgt auch aus einem allgemeinen Satze von St. Banach, Fund. Math. 3 (1922), S. 157. Wir geben, der Vollständigkeit halber, seinen auf den vorliegenden Fall spezialisierten schönen Beweis wieder. Sei also $f_1, f_2, \dots$ schwach gegen f konvergent; es gilt $|f_n| \leq c$, mit festem c , zu beweisen. Es genügt $|(f_n, \varphi)| \leq c|\varphi|$ (für alle φ !) zu zeigen: für $\varphi = f_n$ folgt daraus die Behauptung. Ferner gilt die obige Relation gewiß immer, wenn sie für die φ mit $|\varphi| = \varepsilon$ ($\varepsilon > 0$) gilt, d. h. wenn $|(f_n, \varphi)|$ für alle n und alle $|\varphi| = \varepsilon$ beschränkt ist — also erst recht, wenn dies für alle $|\varphi| \leq \varepsilon$ der Fall ist, d. h. in einer beliebigen Kugel um die 0. Wenn nun alle $|(f_n, \varphi)|$ beschränkt sind, falls φ in einer festen Kugel $\mathfrak{K}$ ($|\varphi - g_0| < \delta$) variiert, sind sie auch in einer Kugel um die 0 beschränkt (nämlich $|\varphi| \leq \delta$): denn ein solches φ ist gewiß Differenz von zwei Punkten von $\mathfrak{K}$ (z. B. $g_0 \pm \frac{1}{2}\varphi$). Wäre also der Satz falsch, so wären die $|(f_n, \varphi)|$ in jeder Kugel $\mathfrak{K}$ unbeschränkt: d. h. es gäbe zu jedem c ein φ_0 in $\mathfrak{K}$ und ein n_0 , so daß $|(f_{n_0}, \varphi_0)| > c$ ist. Aus Stetigkeitsgründen muß dann $|(f_{n_0}, \varphi)| > c$ auch noch in einer geeigneten Kugel um φ_0 gelten, die noch als Teil von $\mathfrak{K}$ gewählt werden kann.

Und nun sei $\mathfrak{K}$ eine beliebige Kugel; $\mathfrak{K}'$ eine Kugel in $\mathfrak{K}$, in der stets $|(f_{n'}, \varphi)| > 1$ ist (n' fest), und zwar sei ihr Radius < 1 ; $\mathfrak{K}''$ eine Kugel in $\mathfrak{K}$, in der stets $|(f_{n''}, \varphi)| > 2$ ist (n'' fest), und zwar sei ihr Radius $< \frac{1}{2}, \dots$. Da $\mathfrak{H}$ vollständig ist (vgl. E., § I, Axiom E) haben alle Kugeln $\mathfrak{K}, \mathfrak{K}', \mathfrak{K}'', \dots$ einen gemeinsamen Punkt $\bar{\varphi}$, und für diesen ist $|(f_{n'}, \bar{\varphi})| > 1, |(f_{n''}, \bar{\varphi})| > 2, \dots$, also $\lim_{n \rightarrow \infty} \sup |(f_n, \bar{\varphi})| = \infty$, entgegen der Annahme.

Wir bemerken noch: Sei $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System. Wenn $f_1, f_2, \dots$ schwach gegen f konvergiert, so ist für jedes $u = 1, 2, \dots$ $\lim_{n \rightarrow \infty} (f_n, \varphi_u) = (f, \varphi_u)$, ferner sind die $|f_n|$ beschränkt. Diese notwendige Bedingung ist auch hinreichend: die Menge der φ mit $\lim_{n \rightarrow \infty} (f_n, \varphi) = (f, \varphi)$ enthält die $\varphi_1, \varphi_2, \dots$, sie ist offenbar eine lin. M., und weil die $|f_n|$ beschränkt (also die (f_n, ψ) in ψ gleichmäßig stetig) sind, ist sie abg. (im starken Sinne) — somit umfaßt sie ganz $\mathfrak{H}$. Damit ist alles bewiesen.

unendlich oft $m = m_0$ sein, d. h. $x_{m_0} = 1$, trotzdem wegen der schwachen Konvergenz nur endlich oft $x_{m_0} = 1$ sein dürfte. —

Zum Schluß sei noch auf den folgenden Umstand hingewiesen. Im Innern einer jeden festen Kugel $\mathfrak{K}$ sind alle f beschränkt, d. h. $|f| \leq c$. Sei $\bar{\varphi}_1, \bar{\varphi}_2, \dots$ ein vollst. norm. orth. System, dann enthält jedes $\mathcal{U}_2(f_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ ein $\mathcal{U}_2(f_0; \bar{\varphi}_1, \dots, \bar{\varphi}_n, \delta)$ (soweit nämlich beide in $\mathfrak{K}$ liegen!): Sei nämlich $\varphi_1 = \sum_{r=1}^{\infty} a_r^1 \bar{\varphi}_r, \dots, \varphi_s = \sum_{r=1}^{\infty} a_r^s \bar{\varphi}_r$, und n mit $\sum_{r=n+1}^{\infty} |a_r^1|^2, \dots, \sum_{r=n+1}^{\infty} |a_r^s|^2 < \frac{\varepsilon^2}{16c^2}$ gewählt, und weiter $\delta \leq \varepsilon : 2 \text{ Max} \left(\sum_{r=1}^n |a_r^1|, \dots, \sum_{r=1}^n |a_r^s| \right)$. Dann gilt für $r = 1, \dots, s$ und alle f des letzteren $\mathcal{U}_2$

$$\begin{aligned} |(f - f_0, \varphi_r)| &\leq \sum_{r=1}^n |a_r^r| |(f - f_0, \bar{\varphi}_r)| + \left(f - f_0, \sum_{r=n+1}^{\infty} a_r^r \bar{\varphi}_r \right) \\ &\leq \sum_{r=1}^n |a_r^r| \delta + |f - f_0| \sqrt{\sum_{r=n+1}^{\infty} |a_r^r|^2} < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon, \end{aligned}$$

d. h. sie gehören zum erstgenannten $\mathcal{U}_2$. Da wir offenbar $\delta = \frac{1}{p}$ ($p = 1, 2, \dots$) wählen können, ist somit das erste Hausdorffsche Abzählbarkeitsaxiom erfüllt³³⁾.

Bekanntlich ist die schwache, nicht aber die starke, Topologie im Innern von $\mathfrak{K}$ sogar kompakt³⁴⁾. Also: die starke Topologie genügt stets dem Abzählbarkeitsaxiom, ist aber weder in ganz $\mathfrak{H}$ noch in $\mathfrak{K}$ kompakt — die schwache hat in ganz $\mathfrak{H}$ keine der beiden Eigenschaften, in $\mathfrak{K}$ dagegen beide.

2. Nach dieser vorhergehenden Orientierung wenden wir uns den (uns eigentlich interessierenden) Verhältnissen in $\mathcal{B}$ zu.

Um die Konvergenz einer O.-Folge $A_1, A_2, \dots$ gegen ein A zu definieren, verlangt man in der Regel entweder die starke Konvergenz aller $A_n f$ gegen Af , oder die schwache — das letztere bedeutet $(A_n f, g) \rightarrow (Af, g)$ für alle f, g . Wir haben also: entweder $|(A_n - A)f| \rightarrow 0$ für alle f , oder $|((A_n - A)f, g)| \rightarrow 0$ für alle f, g . Dies ist die starke bzw. schwache Konvergenz der O. in $\mathcal{B}$. Wie in 1. führen wir diese Konvergenzbegriffe auf entsprechende Topologien in $\mathcal{B}$ zurück. Wir haben dann:

Starke Topologie in $\mathcal{B}$. Sei $\mathcal{U}_3(A_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ die Menge aller A von $\mathcal{B}$ mit $|(A - A_0)\varphi_1| < \varepsilon, \dots, |(A - A_0)\varphi_s| < \varepsilon$; alle $\mathcal{U}_3(A_0; \varphi_1, \dots, \varphi_s, \varepsilon)$

³³⁾ Da es eine stark, also auch schwach, überall dichte Folge in $\mathfrak{H}$ gibt, zeigt man leicht die Gültigkeit auch des zweiten, vgl. Anm. ²⁹⁾.

³⁴⁾ Alle Aussonderungsmethoden (bei Konvergenzbeweisen) im Hilbertschen Raume kommen darauf heraus.

mit beliebigen $\varphi_1, \dots, \varphi_s$ (aus $\mathfrak{H}$, s beliebig) und $\varepsilon > 0$ sind die Umgebungen von A_0 .

Schwache Topologie in $\mathcal{B}$. Sei $\mathcal{U}_4(A_0; \varphi_1, \psi_1, \dots, \varphi_s, \psi_s, \varepsilon)$ die Menge aller A von $\mathcal{B}$ mit $|((A - A_0)\varphi_1, \psi_1)| < \varepsilon, \dots, |((A - A_0)\varphi_s, \psi_s)| < \varepsilon$; alle $\mathcal{U}_4(A_0; \varphi_1, \psi_1, \dots, \varphi_s, \psi_s, \varepsilon)$ mit beliebigen $\varphi_1, \psi_1, \dots, \varphi_s, \psi_s$ (aus $\mathfrak{H}$, s beliebig) und $\varepsilon > 0$ sind die Umgebungen von A_0 .

Vom Erfülltsein der Hausdorffschen Axiome α) bis δ) überzeugt man sich wiederum leicht, ebenso davon, daß diese Topologien auf die oben genannten Konvergenzbegriffe führen. Ferner ist es klar, daß die Operationen αA , $A + B$ in beiden Topologien stetig sind (die letztere in beiden Variablen zugleich). A^* ist in der schwachen stetig (wegen $((A - A_0)\varphi, \psi) = ((A^* - A_0)\varphi, \psi)$), in der starken nicht (wir übergehen die leicht angebbaren Gegenbeispiele). AB ist in beiden Topologien in jeder der beiden Variablen stetig, wenn die andere festgehalten wird: daß sie es stark in A ist, erkennt man, wenn man $\varphi_1, \dots, \varphi_s$ durch $B\varphi_1, \dots, B\varphi_s$ ersetzt; daß sie es stark in B ist, folgt aus der Stetigkeit des (festen) A ; daß sie es schwach in A ist, folgt durch Ersetzen von $\varphi_1, \psi_1, \dots, \varphi_s, \psi_s$ mit $B\varphi_1, \psi_1, \dots, B\varphi_s, \psi_s$; daß sie es schwach in B ist, folgt aus dem vorigen sowie $AB = (B^*A^*)^*$. In beiden Variablen zugleich ist aber A, B in keiner von beiden stetig (auch hier sind Gegenbeispiele so einfach, daß wir sie übergehen).

Da $\mathcal{U}_3(A_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ Teilmenge von $\mathcal{U}_4(A_0; \varphi_1, \psi_1, \dots, \varphi_s, \psi_s, \eta)$ ist (z. B. für $\varepsilon = \eta : \text{Max}(|\psi_1|, \dots, |\psi_s|)$), ist jeder Häufungspunkt bzw. Limes im starken Sinne auch einer im schwachen — die Umkehrung kann schon wegen des verschiedenen Stetigkeitscharakters von A^* nicht gelten.

Das erste Abzählbarkeitsaxiom gilt in keiner dieser Topologien: denn für jede gibt es Mengen mit einem Häufungspunkte, der Limes keiner Teilfolge derselben ist — wir werden beides zeigen, indem wir eine Menge $\mathfrak{A}$ angeben, von der 0 starker Häufungspunkt ist, während keine Teilfolge es auch nur zum schwachen Limes hat.

Zunächst stellen wir aber wieder fest: wenn $A_1, A_2, \dots$ schwach (also erst recht wenn es stark) gegen A konvergiert, so sind die $A_1, A_2, \dots$ gleichmäßig beschränkt — d. h. es existiert ein festes c , so daß für alle n und f $|A_n f| \leq c \cdot |f|$ gilt. Es genügt dies für schwache Konvergenz zu zeigen, wo man es analog zum entsprechenden Satze in § beweist³⁵). Hier-

³⁵) Nach E., Satz 12 β) genügt es $|(A_n f, g)| \leq C \cdot |f| |g|$ zu zeigen. Dies beweisen wir mit der Methode von St. Banach genau so wie die analoge E. Schmidtsche Relation in Anm. ³²), nur daß die Funktion (f_n, φ) (von n und φ) durch die Funktion $(A_n f, g)$ (von n und f, g) zu ersetzen ist — und die Kugeln $\mathfrak{K}, \mathfrak{K}', \mathfrak{K}'', \dots$ (für φ) durch Kugelpaare $\mathfrak{K}, \mathfrak{L}, \mathfrak{K}', \mathfrak{L}', \mathfrak{K}'', \mathfrak{L}'', \dots$ (für f bzw. g).

aus folgt nebenbei, daß AB im Sinne der starken Konvergenz (nicht der Topologie!) in beiden Variablen zugleich stetig ist: aus $A_n \rightarrow A$, $B_n \rightarrow B$ folgt $A_n B_n f \rightarrow ABf$ und wegen $B_n f - Bf \rightarrow 0$ sowie der gleichmäßigen Stetigkeit der A_n $A_n B_n f - A_n Bf \rightarrow 0$ — d. h. $A_n B_n f \rightarrow ABf$ also $A_n B_n \rightarrow AB$ (alle $\rightarrow$ im Sinne der starken Konvergenz von $\mathfrak{S}$ bzw. $\mathcal{B}$). Nun geben wir das gewünschte Gegenbeispiel an.

$\bar{A}_{m,n}$ sei der folgende O. aus $\mathcal{B}$ ($\mathfrak{S}$ sei wieder als Folgenraum der $\{x_1, x_2, \dots\}$ realisiert):

$$\bar{A}_{m,n} \{x_1, x_2, \dots\} = \{y_1, y_2, \dots\},$$

$$y_m = x_m, \quad y_n = m x_n, \quad y_p = 0 \quad \text{für } p \neq m, n;$$

und $\mathfrak{A}$ die Menge aller $\bar{A}_{m,n}$. 0 ist starker Häufungspunkt von $\mathfrak{A}$, d. h. bei gegebenen $\{u_1^1, u_1^2, \dots\}, \dots, \{u_1^s, u_2^s, \dots\}$, ε können wir

$$|\bar{A}_{m,n} \{u_1^1, u_1^2, \dots\}| < \varepsilon, \dots, |\bar{A}_{m,n} \{u_1^s, u_2^s, \dots\}| < \varepsilon,$$

also

$$\sqrt{(u_m^1)^2 + m^2 (u_n^1)^2} < \varepsilon, \dots, \sqrt{(u_m^s)^2 + m^2 (u_n^s)^2} < \varepsilon$$

erzwingen: wir wählen m so, daß $|u_m^1| < \frac{\varepsilon}{\sqrt{2}}, \dots, |u_m^s| < \frac{\varepsilon}{\sqrt{2}}$ gilt (es ist ja $u_p^1 \rightarrow 0, \dots, u_p^s \rightarrow 0$), und dann n so, daß $|u_n^1| < \frac{\varepsilon}{\sqrt{2}m}, \dots, |u_n^s| < \frac{\varepsilon}{\sqrt{2}m}$ gilt. Dagegen konvergiert keine Teilfolge von $\mathfrak{A}$ schwach gegen 0: in einer solchen müßte gleichmäßig $|\bar{A}_{m,n} f| \leq c \cdot |f|$ gelten, d. h. m beschränkt sein, es kämen also nur endlich viele m darin vor — für ein m_0 müßte also unendlich oft $m = m_0$ sein, d. h. $(\bar{A}_{m,n} \varphi, \psi) = 1$, wenn $\varphi = \psi = \{x_1, x_2, \dots\}$ mit $x_{m_0} = 1$, $x_p = 0$ für $p \neq m_0$ ist, trotzdem wegen der schwachen Konvergenz nur endlich oft $(\bar{A}_{m,n} \varphi, \psi) = 1$ sein dürfte.

Der Umstand, daß in der starken Topologie A^* unstetig ist, ist für unsere Zwecke manchmal störend. Wir führen daher die starke Doppelkonvergenz ein: sie besteht, wenn $A_1, A_2, \dots$ gegen A und $A_1^*, A_2^*, \dots$ gegen A^* konvergiert. Aus dem früher Gesagten folgt: der starken Doppelkonvergenz gegenüber sind die Operationen $aA, A^*, A+B, AB$ stetig (die zwei letzten in beiden Variablen zugleich).

Zum Schluß zeigen wir: Wenn die $A_{m,n}$ stark gegen die A_m (m fest, $n \rightarrow \infty$) konvergieren, und die A_n gegen A ($n \rightarrow \infty$), und wenn alle $A_{m,n}$ gleichmäßig stetig sind (d. h. ein festes c mit $|A_{m,n} f| \leq c \cdot |f|$ existiert), so konvergiert eine geeignete Teilfolge A_{M_p, N_p} stark gegen A ($p \rightarrow \infty$). Sei nämlich $f_1, f_2, \dots$ eine in $\mathfrak{S}$ (im starken Sinne, vgl. E., § 1, Axiom C) überall dichte Folge. Wir wählen N_p mit

$$|A_{N_p} f_1 - A f_1| < \frac{1}{p}, \dots, |A_{N_p} f_p - A f_p| < \frac{1}{p}$$

(es ist ja stets $A_n f \rightarrow A f$), und M_p mit

$$|A_{M_p, N_p} f_1 - A_{N_p} f_1| < \frac{1}{p}, \dots, |A_{M_p, N_p} f_p - A_{N_p} f_p| < \frac{1}{p}$$

(es ist ja stets $A_{m,n} f \rightarrow A_n f$) — somit gilt $A_{M_p, N_p} f \rightarrow A f$ für alle $f = f_1, f_2, \dots$, d. h. in einer überall dichten Menge in $\mathfrak{S}$. Da aber alle A_{M_p, N_p} sowie A gleichmäßig stetig sind, gilt es dann für alle f , wie behauptet wurde. Für die starke Doppelkonvergenz gilt offenbar dasselbe.

3. Man kann $\mathcal{B}$ noch auf eine dritte Art topologisieren. Man definiere nämlich die Konvergenz von $A_1, A_2, \dots$ gegen A so: $|(A_n - A)f|$ soll für $n \rightarrow \infty$ und alle f gleichmäßig gegen 0 gehen — d. h. $\leq c_n \cdot |f|$ sein, $c_n \rightarrow 0$ —, oder auch $|((A_n - A)f, g)|$ soll für $n \rightarrow \infty$ und alle f, g es tun — d. h. $\leq c_n \cdot |f| |g|$ sein, $c_n \rightarrow 0$. Nach E. (Satz 12 β)) ist aber beides dasselbe, wir nennen es mit naheliegender Bezeichnung gleichmäßige Konvergenz: wie wir sehen ist es gleichgültig, ob wir die starke oder die schwache so verschärfen. Die zugehörige Topologie ist offenbar diese:

Gleichmäßige Topologie in $\mathcal{B}$. Sei $\mathcal{U}_\varepsilon(A_0; \varepsilon)$ die Menge aller A von $\mathcal{B}$ mit $|(A - A_0)f| < \varepsilon' < \varepsilon$ für alle f und irgendein festes ε' ; alle $\mathcal{U}_\varepsilon(A_0; \varepsilon)$ mit $\varepsilon > 0$ sind die Umgebungen von A_0 .

Daß die Hausdorffschen Axiome α) bis δ) erfüllt sind, ist klar, ebenso das Stimmen des hieraus folgenden Konvergenzbegriffes. Auch sieht man, daß die Operationen $\alpha A, A^*, A + B, AB$ stetig sind (die zwei letzten in beiden Variablen zugleich)³⁶⁾.

Da $\mathcal{U}_\varepsilon(A_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ Teilmenge von $\mathcal{U}_\varepsilon(A_0; \varepsilon)$ ist, ist jeder gleichmäßige Häufungspunkt bzw. Limes auch ein solcher im starken (und somit erst recht im schwachen) Sinne³⁷⁾ — die Umkehrung gilt nicht, da das erste Abzählbarkeitsaxiom im Starken und im Schwachen nicht gilt, wohl aber, wie wir gleich sehen werden, im Gleichmäßigen.

Da wir uns unter den $\mathcal{U}_\varepsilon(A_0; \varepsilon)$ auf die $\mathcal{U}_\varepsilon(A_0; \frac{1}{n})$, $n = 1, 2, \dots$, beschränken können, gilt das erste Abzählbarkeitsaxiom. Übrigens entspringt diese Topologie einer Metrik des linearen Raumes $\mathcal{B}$: wenn wir nämlich

³⁶⁾ A^* darum, weil aus $|Af| \leq c \cdot |f|$ (für alle f) $|A^*f| \leq c \cdot |f|$ folgt — da diese Aussagen mit $|(Af, g)| \leq c \cdot |f| |g|$ bzw. $|(A^*f, g)| \leq c \cdot |f| |g|$ (für alle f, g , vgl. E., Satz 12 β)) gleichwertig sind, und $(A^*f, g) = \overline{(Ag, f)}$ ist.

³⁷⁾ Aus der gleichmäßigen Konvergenz folgt, da $*$ für sie stetig ist, auch die starke Doppelkonvergenz.

als Betrag von A die obere Schranke von $|Af|$ für $|f| = 1$ erklären, d. h. das kleinste c , für welches stets $|Af| \leq c \cdot |f|$ ist — er heiße $|A|$ —, dann gelten erstens offenbar die folgenden Relationen:

$$|aA| = |a| |A|, \quad |A+B| \leq |A| + |B|, \quad |AB| \leq |A| |B|,$$

also ist $|A-B|$ eine Entfernung im linearen Raume $\mathcal{B}$ (vgl. E., § I, Axiom B, ferner Definition 4 und Anm. ²⁷), und zweitens ist $\mathcal{U}_5(A_0, \varepsilon)$ die Menge aller A mit $|A-A_0| < \varepsilon$: d. h. die gleichmäßige Topologie entsteht aus dieser Metrik. —

Wie in § so ist auch in $\mathcal{B}$ innerhalb einer jeden festen Kugel $\mathcal{K}$ (wo alle A $|A| \leq c$ genügen) der Charakter der starken und der schwachen Topologie wesentlich verändert: hier erfüllen beide das erste Abzählbarkeitsaxiom. Sei nämlich $\bar{\varphi}_1, \bar{\varphi}_2, \dots$ ein vollst. norm. orth. System, dann enthält jedes $\mathcal{U}_4(A_0; \varphi_1, \dots, \varphi_s, \varepsilon)$ bzw. $\mathcal{U}_5(A_0; \varphi_1, \psi_1, \dots, \varphi_s, \psi_s, \varepsilon)$ ein $\mathcal{U}_4(A_0; \bar{\varphi}_1, \dots, \bar{\varphi}_n, \delta)$ bzw. $\mathcal{U}_5(A_0; \bar{\varphi}_{k_1}, \bar{\varphi}_{l_1}, \dots, \bar{\varphi}_{k_m}, \bar{\varphi}_{l_m}, \delta)$: Wir setzen im ersten Falle $\varphi_1 = \sum_{r=1}^{\infty} a_r^1 \bar{\varphi}_r, \dots, \varphi_s = \sum_{r=1}^{\infty} a_r^s \bar{\varphi}_r$, und wählen n mit

$$\sum_{r=n+1}^{\infty} |a_r^1|^2 < \frac{\varepsilon^2}{16c^2}, \dots, \sum_{r=n+1}^{\infty} |a_r^s|^2 < \frac{\varepsilon^2}{16c^2}$$

sowie

$$\delta \leq \varepsilon : 2 \text{Max} \left(\sum_{r=1}^{\infty} |a_r^1|, \dots, \sum_{r=1}^{\infty} |a_r^s| \right).$$

Dann gilt für $r=1, \dots, s$ und alle A des letzteren $\mathcal{U}_4$

$$\begin{aligned} |(A-A_0)\varphi_r| &\leq \sum_{r=1}^n |a_r^r| |(A-A_0)\bar{\varphi}_r| + |(A-A_0) \sum_{r=n+1}^{\infty} a_r^r \bar{\varphi}_r| \\ &\leq \sum_{r=1}^n |a_r^r| \delta + 2c \sqrt{\sum_{r=n+1}^{\infty} |a_r^r|^2} < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon, \end{aligned}$$

d. h. sie gehören zum erstgenannten $\mathcal{U}_4$. Im zweiten Falle setzen wir $\varphi_1 = \sum_{r=1}^{\infty} a_r^1 \bar{\varphi}_r, \psi_1 = \sum_{r=1}^{\infty} b_r^1 \bar{\varphi}_r, \dots, \varphi_s = \sum_{r=1}^{\infty} a_r^s \bar{\varphi}_r, \psi_s = \sum_{r=1}^{\infty} b_r^s \bar{\varphi}_r$, und wählen n mit $\sum_{r=n+1}^{\infty} |a_r^1|^2 < \eta^2, \sum_{r=n+1}^{\infty} |b_r^1|^2 < \eta^2, \dots, \sum_{r=n+1}^{\infty} |a_r^s|^2 < \eta^2, \sum_{r=n+1}^{\infty} |b_r^s|^2 < \eta^2$, wo $\eta > 0$ mit

$$4c \text{Max} (|\varphi_1|, |\psi_1|, \dots, |\varphi_s|, |\psi_s|) \eta + \eta^2 = \frac{\varepsilon}{2}$$

gewählt ist, und setzen ferner $\delta \leq \varepsilon : 2 \text{Max} \left(\left(\sum_{r=1}^n |a_r^1| \right)^2, \dots, \left(\sum_{r=1}^n |a_r^s| \right)^2 \right)$. Schließlich sei $m = n^2$ und $k_1, l_1, \dots, k_m, l_m$ alle Paare ϱ, σ $\varrho \leq n, \sigma \leq n$. Dann gilt für $r=1, \dots, s$ und alle A des letzteren $\mathcal{U}_5$

$$\begin{aligned}
& \left| ((A - A_0) \varphi_r, \psi_r) - \left((A - A_0) \sum_{v=1}^n a_v^r \bar{\varphi}_v, \sum_{v=1}^n b_v^r \bar{\varphi}_v \right) \right| \\
&= \left| - \left((A - A_0) \varphi_r, \sum_{v=n+1}^{\infty} b_v^r \bar{\varphi}_v \right) - \left((A - A_0) \sum_{v=n+1}^{\infty} a_v^r \bar{\varphi}_v, \psi_1 \right) \right. \\
&\quad \left. + \left((A - A_0) \sum_{v=n+1}^{\infty} a_v^r \bar{\varphi}_v, \sum_{v=n+1}^{\infty} b_v^r \bar{\varphi}_v \right) \right| \\
&\leq 2c \cdot |\varphi_r| \cdot \sqrt{\sum_{v=n+1}^{\infty} |b_v^r|^2} + 2c \cdot \sqrt{\sum_{v=n+1}^{\infty} |a_v^r|^2} \cdot |\psi_r| \\
&\quad + \sqrt{\sum_{v=n+1}^{\infty} |a_v^r|^2} \sqrt{\sum_{v=n+1}^{\infty} |b_v^r|^2} \\
&< 4c \operatorname{Max} (|\varphi_r|, |\psi_r|) \eta + \eta^2 \leq \frac{\varepsilon}{2},
\end{aligned}$$

$$\begin{aligned}
\left| \left((A - A_0) \sum_{v=1}^n a_v^r \bar{\varphi}_v, \sum_{v=1}^n b_v^r \bar{\varphi}_v \right) \right| &\leq \sum_{\mu, v=1}^n |a_{\mu}^r| |b_v^r| \cdot |(A - A_0) \bar{\varphi}_{\mu}, \bar{\varphi}_v| \\
&\leq \sum_{\mu, v=1}^n |a_{\mu}^r| |b_v^r| \cdot \delta \leq \frac{\varepsilon}{2},
\end{aligned}$$

also

$$|((A - A_0) \varphi_r, \psi_r)| \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon,$$

d. h. sie gehören zum erstgenannten $\mathcal{U}_5$. Da wir offenbar $\delta = \frac{1}{p}$ ($p = 1, 2, \dots$) wählen können, ist damit die obige Behauptung erwiesen.

4. Zum Schluß bemerken wir: jede Teilmenge $\mathcal{M}$ von $\mathcal{B}$ hat eine abzählbare Teilmenge $\mathcal{N}$, so daß jedes Element A von $\mathcal{M}$ Limes einer stark doppelkonvergenten Teilfolge von $\mathcal{N}$ ist³⁵⁾.

Sei $\mathcal{M}_c$ die Menge aller A von $\mathcal{M}$ mit $|A| < c$: es genügt die Behauptung für alle $\mathcal{M}_1, \mathcal{M}_2, \dots$ zu beweisen (sei dann $\mathcal{N}$ bzw. gleich $\mathcal{N}_1, \mathcal{N}_2, \dots$), da wir $\mathcal{N} = \mathcal{N}_1 + \mathcal{N}_2 + \dots$ setzen können — somit dürfen wir von Anfang an voraussetzen, daß in $\mathcal{M}$ durchweg $|A| < c$ ist.

Wir realisieren $\mathfrak{S}$ als Folgenraum (der $\{x_1, x_2, \dots\}$ mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$), dann hat A eine Matrix $\{a_{\mu, v}\}$, indem allgemein

³⁵⁾ Also ist $\mathcal{N}$ in $\mathcal{M}$ stark wie schwach überall dicht, wir können insbesondere $\mathcal{M} = \mathcal{B}$ setzen. Für die gleichmäßige Topologie ist dies gewiß falsch, da dort leicht eine Menge von Kontinuum vielen A angegeben werden kann, die paarweise die Entfernung 1 haben. — Wenn wir $\mathcal{M}$ einer Kugel $\mathcal{K}$ gleichsetzen, so können wir in $\mathcal{K}$ aus der Gültigkeit des ersten Abzählbarkeitsaxioms (im starken wie im schwachen Sinne) leicht auch die des zweiten folgern (vgl. Anm. ²⁹⁾, ³³⁾). Wir erwähnen noch, daß die schwache Topologie in $\mathcal{K}$ kompakt ist (vgl. Ende von 1. und Anm. ³⁴⁾).

$$A\{x_1, x_2, \dots\} = \{y_1, y_2, \dots\},$$

$$y_\nu = \sum_{\mu=1}^{\infty} a_{\mu,\nu} x_\mu$$

(A ist ja lin. und stetig!) gilt. Sei nun $p = 1, 2, \dots$ und $a_{\mu,\nu}^{(p)}$ ($\mu, \nu = 1, \dots, p$) irgendwelche rationalen Zahlen, so daß für alle $\mu, \nu = 1, \dots, p$ $|a_{\mu,\nu}^{(p)} - a_{\mu,\nu}| < \frac{1}{p}$ ist, wir definieren $A^{(p)}$ durch

$$A^{(p)}\{x_1, x_2, \dots\} = \{y_1, y_2, \dots\},$$

$$y_\nu = \begin{cases} \sum_{\mu=1}^p a_{\mu,\nu}^{(p)} x_\mu, & \text{für } \nu = 1, \dots, p, \\ 0, & \text{für die übrigen } \nu. \end{cases}$$

Für $p \geq m$ haben wir also:

$$\begin{aligned} |A^{(p)}\varphi_m - A\varphi_m|^2 &= \left| \sum_{\mu=1}^p a_{m\mu}^{(p)} \varphi_\mu - \sum_{\mu=1}^{\infty} a_{m\mu} \varphi_\mu \right|^2 \\ &= \left| \sum_{\mu=1}^p (a_{m\mu}^{(p)} - a_{m\mu}) \varphi_\mu - \sum_{\mu=p+1}^{\infty} a_{m\mu} \varphi_\mu \right|^2 \\ &= \sum_{\mu=1}^p |a_{m\mu}^{(p)} - a_{m\mu}|^2 + \sum_{\mu=p+1}^{\infty} |a_{m\mu}|^2 < \frac{1}{p} + \sum_{\mu=p+1}^{\infty} |a_{m\mu}|^2. \end{aligned}$$

Für $p \rightarrow \infty$ strebt dies offenbar gegen 0,³⁰⁾ und dies gilt für alle $m = 1, 2, \dots$. Wir haben also $A^{(p)}\varphi_m \rightarrow A\varphi_m$; ebenso zeigt man $A^{(p)*}\varphi_m \rightarrow A^*\varphi_m$.

Nun betrachten wir alle $O. B$ aus $\mathcal{B}$, deren Matrix $\{b_{\mu,\nu}\}$ aus lauter rationalen Zahlen besteht, von denen überdies nur endlich viele $\neq 0$ sind — sie bilden offenbar eine abzählbare Menge $B_1, B_2, \dots$. Die $A^{(p)}$ von vorhin sind solche B , wir haben also zu jedem A von $\mathcal{B}$ eine Folge $B_{r_1}, B_{r_2}, \dots$, so daß $B_{r_p}f \rightarrow Af$, $B_{r_p}^*f \rightarrow A^*f$ für alle $f = \varphi_1, \varphi_2, \dots$.

Jetzt gehen wir zu $\mathcal{M}$ über. Wir untersuchen für jedes Paar r, p ($= 1, 2, \dots$) ob es ein A von $\mathcal{M}$ mit $|B_{r_p}f - Af| < \frac{1}{p}$, $|B_{r_p}^*f - A^*f| < \frac{1}{p}$, für $f = \varphi_1, \dots, \varphi_p$, gibt — wenn es welche gibt, so wählen wir ein solches A aus und nennen es $A_{r,p}$. Die $A_{r,p}$ bilden eine abzählbare Teilmenge $\mathcal{N}$ von $\mathcal{M}$; es soll gezeigt werden, daß sie das Gewünschte leistet.

³⁰⁾ Weil $\sum_{\mu=1}^{\infty} |a_{m\mu}|^2$ konvergiert. Die Konvergenz aller Reihen $\sum_{m=1}^{\infty} |a_{mn}|^2$, $\sum_{n=1}^{\infty} |a_{mn}|^2$ sieht man so ein: Die zweite konvergiert, weil $A\{0, \dots, 0, \frac{1}{m}, 0, \dots\} = \{a_{m1}, a_{m2}, \dots\}$ ist, die erste aber infolgedessen auch, wenn man statt A A^* betrachtet (wobei aus a_{nm} $\overline{a_{mn}}$ wird).

Sei A irgendein Element von $\mathcal{B}$, $p = 1, 2, \dots$. Wir bilden die vorhin genannte Folge $B_{r_1}, B_{r_2}, \dots$; für diese streben alle $|B_{r_q}f - Af|, |B_{r_q}^*f - A^*f|$ mit $q \rightarrow \infty$ gegen 0 ($f = \varphi_1, \varphi_2, \dots$) — wenn wir also $n = r_q$ groß genug wählen, so ist $|B_n f - Af| < \frac{1}{p}$, $|B_n^* f - A^* f| < \frac{1}{p}$ für alle $f = \varphi_1, \dots, \varphi_p$. Daher kann $A_{n,p}$ gebildet werden (weil z. B. A so ist), und für alle $f = \varphi_1, \dots, \varphi_p$ gilt $|B_n f - A_{n,p} f| < \frac{1}{p}$, $|B_n^* f - A_{n,p}^* f| < \frac{1}{p}$, also $|A_{n,p} f - Af| < \frac{2}{p}$, $|A_{n,p}^* f - A^* f| < \frac{2}{p}$. Wir nennen $A_{n,p}$ (auch n hängt von p ab!) kurz $\bar{A}_p$, dann gilt offenbar $\bar{A}_p f \rightarrow Af$, $\bar{A}_p^* f \rightarrow A^* f$ für alle $f = \varphi_1, \varphi_2, \dots$. Also auch für alle Linearaggregate endlich vieler derselben, d. h. in einer in § (stark) überall dichten Menge; und somit, da die $\bar{A}_p$ alle zu $\mathcal{N}$, also zu $\mathcal{M}$ gehören und darum (in p) gleichmäßig stetig sind, für alle f überhaupt. Die $\bar{A}_1, \bar{A}_2, \dots$ sind somit stark doppelkonvergent gegen A — sie bilden aber eine Teilfolge von $\mathcal{N}$.

II. Eigenschaften von $\mathbf{R}(\mathcal{M})$ und $\mathcal{M}'$.

1. Wir gehen nun an die Definition der schon in Einleitung 2. erwähnten Operationen $\mathbf{R}(\mathcal{M})$ und $\mathcal{M}'$ heran.

Definition 1. Eine Teilmenge $\mathcal{M}$ von $\mathcal{B}$ heißt ein Ring, wenn sie erstens mit A, B auch $aA, A^*, A+B, AB$ enthält, und zweitens schwach abg. ist⁴⁰⁾.

Definition 2. Da der Durchschnitt beliebig vieler Ringe wieder ein solcher ist, ist insbesondere der Durchschnitt aller, eine gegebene Teilmenge $\mathcal{M}$ von $\mathcal{B}$ umfassender Ringe ein Ring — der kleinste, der $\mathcal{M}$ umfaßt. Er heie $\mathbf{R}(\mathcal{M})$, der von $\mathcal{M}$ erzeugte Ring.

Definition 3. Sei $\mathcal{M}$ eine Teilmenge von $\mathcal{B}$. Die Menge aller A von $\mathcal{B}$, für die A und A^* mit jedem B von $\mathcal{M}$ vert. ist (vgl. Anm. ¹⁶⁾), heie $\mathcal{M}'$. So können auch Iterierte $\mathcal{M}'', \mathcal{M}''', \dots$ gebildet werden.

$\mathcal{M}'$ ist stets ein Ring: da es mit A, B auch $aA, A^*, A+B, AB$ enthält, ist klar, aber es ist auch schwach abg. — denn wenn A ein schwacher Häufungspunkt von $\mathcal{M}'$ ist und B zu $\mathcal{M}$ gehört, so folgt daraus, da für alle A' von $\mathcal{M}'$ $A'B = BA'$, $A'^*B = BA'^*$ gilt und diese vier Ausdrücke alle in A (schwach, bei festem B) stetig sind, $AB = BA$, $A^*B = BA^*$.

Wenn A zu $\mathcal{M}$ und B zu $\mathcal{M}'$ gehört, so sind B, B^* mit A vert., also A, A^* mit B : daher gehört A zu $\mathcal{M}''$. Also ist $\mathcal{M} \subset \mathcal{M}''$. Ferner ist es klar, da aus $\mathcal{M} \subset \mathcal{N}$ $\mathcal{M}' \supset \mathcal{N}'$ folgt. Wenn wir dies auf $\mathcal{M} \subset \mathcal{M}''$

⁴⁰⁾ Man beachte: schwach abg. ist mehr wie stark abg.!

anwenden, wird $\mathcal{M}' > \mathcal{M}'''$; wenn wir aber darin bloß $\mathcal{M}$ durch $\mathcal{M}'$ ersetzen, so wird $\mathcal{M}' < \mathcal{M}'''$ — also ist $\mathcal{M}' = \mathcal{M}'''$. Ersetzen wir hierin $\mathcal{M}$ durch $\mathcal{M}', \mathcal{M}'', \dots$, so wird:

$$\mathcal{M}' = \mathcal{M}''' = \mathcal{M}^V = \dots, \quad \mathcal{M}'' = \mathcal{M}^{IV} = \mathcal{M}^{VI} = \dots$$

Da $\mathcal{M}'$ stets ein Ring ist und offenbar die 1 enthält, gilt dasselbe für $\mathcal{M}''$; wegen $\mathcal{M} < \mathcal{M}''$ ist also $(\mathcal{M}, 1)^{41)} < \mathcal{M}''$, d. h. $R(\mathcal{M}, 1) < \mathcal{M}''$ (also auch $R(\mathcal{M}) < \mathcal{M}''$). Wir werden übrigens bald zeigen (Satz 5), daß sogar $R(\mathcal{M}, 1) = \mathcal{M}''$ gilt.

Wir nennen eine Menge $\mathcal{M}$ Abelsch, wenn alle Elemente von ihr sowie deren $*$ miteinander vert. sind — für Mengen, die mit A auch A^* enthalten, z. B. Ringe, genügt also die Vert. aller Elemente. Abelsch sein bedeutet offenbar $\mathcal{M} < \mathcal{M}'$; da hieraus $\mathcal{M}'' < \mathcal{M}'''$ folgt, ist mit $\mathcal{M}$ auch $\mathcal{M}''$ Abelsch; und wegen $\mathcal{M} < \mathcal{M}''$ gilt auch die Umkehrung. Aus $\mathcal{M} < R(\mathcal{M}) < \mathcal{M}''$ folgt, daß auch $R(\mathcal{M})$ mit diesen Mengen gleichzeitig Abelsch ist — also: der Abelsche Charakter ist für alle drei Mengen $\mathcal{M}, R(\mathcal{M}), \mathcal{M}''$ gleichwertig.

Sei $\mathcal{M}$ ein Ring. Wenn er Abelsch ist, ist für jedes A von $\mathcal{M}$ A mit A^* vert., d. h. A normal (vgl. Einleitung 2., bei Anm.¹⁶⁾). Wenn er nicht Abelsch ist, so enthält er zwei unvert. A, B ; wir setzen $A_1 = \frac{A+A^*}{2}$, $A_2 = \frac{A-A^*}{2i}$, $B_1 = \frac{B+B^*}{2}$, $B_2 = \frac{B-B^*}{2i}$ — die A_1, A_2, B_1, B_2 sind H. O. aus $\mathcal{M}$ und gewiß nicht alle vert. (wegen $A = A_1 + iA_2$, $B = B_1 + iB_2$), also können wir A, B gleich als H. O. annehmen. Dann gehört $C = A + iB$ zu $\mathcal{M}$ und ist mit $C^* = A - iB$ unvert., also nicht normal. Also: Ein Ring $\mathcal{M}$ ist dann und nur dann Abelsch, wenn alle seine Elemente normal sind — ein beliebiges $\mathcal{M}$ also, wenn dies für $R(\mathcal{M})$ gilt.

2. Jetzt sollen die Beziehungen der obigen Begriffsbildungen zu den P. O. und unitären O. erörtert werden.

Satz 1. $\mathcal{M}$ sei ein Ring, A ein (beschr.) H. O., $E(\lambda)$ seine Z. d. E.⁴²⁾.

⁴¹⁾ Die Vereinigungsmenge von $\mathcal{M}$ mit der 1.

⁴²⁾ In E. (Definition 17) wurde eine Z. d. E. als eine Schar von P. O. $E(\lambda)$ ($-\infty < \lambda < \infty$) mit den folgenden Eigenschaften definiert:

α) Wenn $\lambda \leq \mu$ ist, so ist $E(\lambda)$ Teil von $E(\mu)$.

β) Wenn $\lambda \geq \lambda_0$, $\lambda \rightarrow \lambda_0$, so gilt für jedes f $E(\lambda)f \rightarrow E(\lambda_0)f$.

γ) Wenn $\lambda \rightarrow -\infty$ oder $\rightarrow +\infty$, so ist stets $E(\lambda)f \rightarrow 0$ bzw. $\rightarrow f$.

[Alle Konvergenzen, wie stets in E., stark. Wir würden jetzt für β) sagen: $E(\lambda)$ ist in der starken Topologie von $\mathcal{B}$ eine nach rechts halbstetige Funktion von λ ; und für γ): es kann für $\lambda = \pm \infty$, in der starken Topologie von $\mathcal{B}$ stetig als 1 bzw. 0 definiert werden.] Die Z. d. E. $E(\lambda)$ gehört zum (nicht notwendig beschr.) H. O. A (vgl. E. Satz 36), wenn weiter gilt:

A gehört dann und nur dann zu $\mathcal{M}$, wenn alle $E(\lambda)$, $\lambda < 0$, und $1 - E(\lambda)$, $\lambda \geq 0$, zu $\mathcal{M}$ gehören. Wenn $\mathcal{M}$ die 1 enthält, können wir einfacher sagen: wenn alle $E(\lambda)$ zu $\mathcal{M}$ gehören.

Beweis. Die zweite Behauptung folgt aus der ersten; betrachten wir also diese.

Zunächst ist die Bedingung hinreichend. Es mögen nämlich alle $E(\lambda)$, $\lambda < 0$, und $1 - E(\lambda)$, $\lambda \geq 0$, zu $\mathcal{M}$ gehören, und $-c \leq \lambda \leq c$ sei das in Anm. ⁴²⁾ erwähnte Intervall, außerhalb dessen $E(\lambda)$ gleich 0 bzw. 1 ist. Sei K_ε eine Kette von Zahlen von $-c$ bis c : $-c = \lambda_0 < \lambda_1 < \dots < \lambda_{n-1} < \lambda_n = c$, mit einer Maschenweite $\leq \varepsilon$, d. h. $\lambda_\nu - \lambda_{\nu-1} \leq \varepsilon$ für $\nu = 1, 2, \dots, n$ ($\varepsilon > 0$), und in der die 0 vorkommt — etwa $\lambda_m = 0$. Dann gehört

$$\begin{aligned} A^{K_\varepsilon} &= \sum_{\nu=1}^n \lambda_\nu (E(\lambda_\nu) - E(\lambda_{\nu-1})) \\ &= \sum_{\nu=1}^{m-1} \lambda_\nu (E(\lambda_\nu) - E(\lambda_{\nu-1})) + \sum_{\nu=m+1}^n \lambda_\nu ((1 - E(\lambda_{\nu-1})) - (1 - E(\lambda_\nu))) \end{aligned}$$

auch zu $\mathcal{M}$, und es ist

$$(A^{K_\varepsilon} f, g) = \sum_{\nu=1}^n \lambda_\nu ((E(\lambda_\nu) f, g) - (E(\lambda_{\nu-1}) f, g)) \dots$$

Spektralform. Af hat dann und nur dann Sinn, wenn

$$\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2$$

endlich ist (Stieltjessches Integral, da der Integrand stets ≥ 0 ist, und der Ausdruck hinter dem d , wie sich zeigen läßt, nie abnimmt, ist das Integral endlich [konvergent] oder $+\infty$ [eigentlich divergent] — nun soll das erstere zutreffen). Ist dies der Fall, so gilt für alle g

$$(Af, g) = \int_{-\infty}^{\infty} \lambda d(E(\lambda)f, g)$$

(das Integral konvergiert absolut).

In E. wurde festgestellt, daß die sogenannten hypermax. H. O. solche Z. d. E. haben, und nur diese — dabei entsprechen sich alle hypermax. H. O. und alle Z. d. E. eindeutig. Dies war die Verallgemeinerung der Hilbertschen Spektralform ins Unbeschr. (vgl. E. Satz 36).

Für beschr. H. O. A ist die Z. d. E. stets vorhanden (dies ist das Resultat der Hilbertschen Theorie, vgl. auch E., wo sich zeigt, daß alle beschr. A hypermax. sind, Satz 48, α), β). Für die Z. d. E. dieser H. O. ist es aber charakteristisch, daß sich die Eigenschaft γ) dahin verschärfen läßt, daß $E(\lambda)$ für genügend kleine bzw. große λ überhaupt konstant ist — also gleich 0 bzw. 1. Wenn dies etwa für $\lambda \leq -c$ bzw. $\lambda \geq c$ eintritt, so können wir also alle $\int_{-\infty}^{\infty}$ durch $\int_{-c}^c$ ersetzen (und γ) ist von selbst erfüllt).

Diese Kennzeichnung der Beschr. ist eines der Resultate der Hilbertschen Theorie, wir werden sie aber — lediglich der Vollständigkeit halber — im Anhang I kurz begründen.

Mit $\varepsilon \rightarrow 0$ strebt dies für alle f, g gegen $\int_{-c}^c \lambda d(E(\lambda)f, g) = (Af, g)$, d. h. A ist schwacher Limes der A^{K_ε} .⁴³⁾ Also auch schwacher Häufungspunkt von $\mathcal{M}$ — somit gehört es zu $\mathcal{M}$.

Weiter ist die Bedingung notwendig. Es gehöre jetzt A zu $\mathcal{M}$; dann gehören auch $A, A^2, A^3, \dots$ dazu, und alle ihre Linearaggregate — d. h. alle $p(A)$, wo $p(x)$ jedes Polynom mit $p(0) = 0$ sein kann. Man beweist leicht:

$$(p(A)f, g) = \int_{-c}^c p(\lambda) d(E(\lambda)f, g). \quad {}^{44)}$$

Für $\lambda_0 < 0$ wählen wir ein $\varepsilon > 0$ und $p_\varepsilon(x)$ den Bedingungen

$$\left. \begin{aligned} 1 - \varepsilon < p_\varepsilon(x) < 1 + \varepsilon & \text{ für } -c \leq x \leq \lambda_0 \\ -\varepsilon < p_\varepsilon(x) < 1 + \varepsilon & \text{ für } \lambda_0 \leq x \leq \lambda_0 + \varepsilon, \\ -\varepsilon < p_\varepsilon(x) < \varepsilon & \text{ für } \lambda_0 + \varepsilon \leq x \leq c, \end{aligned} \right\} \text{ sowie } p_\varepsilon(0) = 0$$

entsprechend (man erkennt mühelos, daß dies möglich ist). Dann ist für jedes f und g

$$\begin{aligned} (p_\varepsilon(A)f - E(\lambda_0)f, g) &= \int_{-c}^c p_\varepsilon(\lambda) d(E(\lambda)f, g) - \int_{-c}^{\lambda_0} d(E(\lambda)f, g) \\ &= \int_{-c}^{\lambda_0} (p_\varepsilon(\lambda) - 1) d(E(\lambda)f, g) + \int_{\lambda_0}^{\lambda_0 + \varepsilon} p_\varepsilon(\lambda) d(E(\lambda)f, g) + \int_{\lambda_0 + \varepsilon}^c p_\varepsilon(\lambda) d(E(\lambda)f, g). \end{aligned}$$

In diesen drei Integralen sind die Integranden bzw. absolut $\leq \varepsilon, 1 + \varepsilon, \varepsilon$, wir können daher mit Hilfe der Abschätzungsmethode aus Anm. ⁴³⁾

⁴³⁾ Es ist sogar gleichmäßiger Limes derselben. Denn wenn $f(\lambda)$ durch $f_\varepsilon(\lambda) = \lambda_\nu$ für $\lambda_{\nu-1} < \lambda \leq \lambda_\nu$ definiert wird, so ist

$$\begin{aligned} (A^{K_\varepsilon}f, g) &= \int_{-c}^c f_\varepsilon(\lambda) d(E(\lambda)f, g), \quad (Af, g) = \int_{-c}^c \lambda d(E(\lambda)f, g), \\ ((A^{K_\varepsilon} - A)f, g) &= \int_{-c}^c (f_\varepsilon(\lambda) - \lambda) d(E(\lambda)f, g). \end{aligned}$$

Nun ist stets $|f_\varepsilon(\lambda) - \lambda| \leq \varepsilon$, hieraus folgt wegen

$$|(E(\mu)f, g) - (E(\lambda)f, g)| \leq \sqrt{|E(\mu)f|^2 - |E(\lambda)f|^2} \sqrt{|E(\mu)g|^2 - |E(\lambda)g|^2}$$

(vgl. die in E. beim Beweise von Satz 36) durchgeführte Abschätzung)

$$|((A^{K_\varepsilon} - A)f, g)| \leq \varepsilon \cdot |f| \cdot |g|,$$

also (vgl. § I. 3.) $|A^{K_\varepsilon} - A| \leq \varepsilon$. Damit ist alles bewiesen.

⁴⁴⁾ Die Rechnung entspricht in allen Einzelheiten der in E. Anhang I (beim Beweis von Satz 14*) für unitäre O. ausgeführten. Der Kunstgriff geht auf F. Riesz zurück.

$$\begin{aligned}
& |(p_\varepsilon(A)f - E(\lambda_0)f, g)| \\
& \leq \varepsilon \cdot |E(\lambda_0)f| |E(\lambda_0)g| + (1+\varepsilon) \cdot \sqrt{|E(\lambda_0+\varepsilon)f|^2 - |E(\lambda_0)f|^2} \sqrt{|E(\lambda_0+\varepsilon)g|^2 - |E(\lambda_0)g|^2} \\
& \quad + \varepsilon \cdot \sqrt{|f|^2 - |E(\lambda_0+\varepsilon)f|^2} \sqrt{|g|^2 - |E(\lambda_0+\varepsilon)g|^2}^{45)} \\
& \leq \varepsilon \cdot |f| |g| + \sqrt{|E(\lambda_0+\varepsilon)f|^2 - |E(\lambda_0)f|^2} \sqrt{|E(\lambda_0+\varepsilon)g|^2 - |E(\lambda_0)g|^2} \\
& \leq \varepsilon \cdot |f| |g| + \sqrt{|E(\lambda_0+\varepsilon)f|^2 - |E(\lambda_0)f|^2} |g|
\end{aligned}$$

schließen. Da dies für alle g gilt, so folgt daraus ⁴⁶⁾

$$|p_\varepsilon(A)f - E(\lambda_0)f| \leq \varepsilon \cdot |f| + \sqrt{|E(\lambda_0+\varepsilon)f|^2 - |E(\lambda_0)f|^2},$$

was mit $\varepsilon \rightarrow 0$ gegen 0 strebt. Somit ist für $\varepsilon \rightarrow 0$ (etwa in der Folge $\varepsilon = 1, \frac{1}{2}, \frac{1}{3}, \dots$) $p_\varepsilon(A)f \rightarrow E(\lambda_0)f$ (stark in $\mathfrak{H}$!) für jedes f — somit konvergieren die $p_\varepsilon(A)$ stark (in $\mathcal{B}$) gegen $E(\lambda_0)$. Also gehört $E(\lambda_0)$ zu $\mathcal{M}$.

Für $\lambda \geq 0$ wählen wir Polynome $q_\varepsilon(x)$ mit

$$\left. \begin{aligned}
- \varepsilon < q_\varepsilon(x) < \varepsilon & \quad \text{für} \quad -c \leq x \leq \lambda_0, \\
- \varepsilon < q_\varepsilon(x) < 1 + \varepsilon & \quad \text{für} \quad \lambda_0 \leq x \leq \lambda_0 + \varepsilon, \\
1 - \varepsilon < q_\varepsilon(x) < 1 + \varepsilon & \quad \text{für} \quad \lambda_0 + \varepsilon \leq x \leq c,
\end{aligned} \right\} \text{ sowie } q_\varepsilon(0) = 0$$

(auch dies ist ohne weiteres möglich), und erkennen dann durch ganz analoge Betrachtungen wie vorhin, daß die $q_\varepsilon(A)$ für $\varepsilon \rightarrow 0$ stark (in $\mathcal{B}$) gegen $1 - E(\lambda_0)$ konvergieren. In diesem Falle gehört also $1 - E(\lambda_0)$ zu $\mathcal{M}$.

Damit ist alles bewiesen.

Satz 2. Sei $\mathcal{M}$ ein Ring, $\mathcal{M}^P, \mathcal{M}^U$ die Menge aller P.O. bzw. aller unitären O. aus $\mathcal{M}$. Es ist stets $R(\mathcal{M}^P) = \mathcal{M}$; $\mathcal{M}^U$ ist leer, wenn 1 nicht zu $\mathcal{M}$ gehört, im entgegengesetzten Falle aber ist $R(\mathcal{M}^U) = \mathcal{M}$.

Beweis. Nach Satz 1 sind zwei Ringe $\mathcal{M}, \mathcal{N}$ identisch, wenn $\mathcal{M}^P = \mathcal{N}^P$ ist — denn sie enthalten dieselben H.O., also überhaupt dieselben O., da für die Zugehörigkeit eines beliebigen O. A die der zwei H.O. $\frac{A+A^*}{2}, \frac{A-A^*}{2i}$ charakteristisch ist. Nun ist $\mathcal{M}^P \subset \mathcal{M}, R(\mathcal{M}^P) \subset R(\mathcal{M}) = \mathcal{M}$, $(R(\mathcal{M}^P))^P \subset \mathcal{M}^P$, und andererseits $R(\mathcal{M}^P) \supset \mathcal{M}^P, (R(\mathcal{M}^P))^P \supset \mathcal{M}^P$. Also gilt $(R(\mathcal{M}^P))^P = \mathcal{M}^P$, d. h. $R(\mathcal{M}^P) = \mathcal{M}$.

⁴⁵⁾ Wie in E., Beweis von Satz 36 (vgl. die vorhin zitierte Anm. ⁴³⁾) wird hier die Ungleichheit

$$\sqrt{a_1} \sqrt{b_1} + \dots + \sqrt{a_p} \sqrt{b_p} \leq \sqrt{(a_1 + \dots + a_p)(b_1 + \dots + b_p)}$$

angewandt.

⁴⁶⁾ Man setze $g = p_\varepsilon(A)f - E(\lambda_0)f$.

Wenn U ein unitäres Element von $\mathcal{M}$ ist, so gehören auch U^* und $UU^*=1$ zu $\mathcal{M}$; wenn also 1 nicht dazugehört, so ist $\mathcal{M}^U$ leer. Gehöre nun 1 zu $\mathcal{M}$. Erstens ist $\mathcal{M}^U \subset \mathcal{M}$, $R(\mathcal{M}^U) \subset R(\mathcal{M}) = \mathcal{M}$. Zweitens gehöre E zu $\mathcal{M}^P$, dann gehören $E, 1$ und somit $2E-1$ zu $\mathcal{M}$; dieses ist aber unitär, weil E ein P.O. ist⁴⁷⁾, gehört also zu $\mathcal{M}^U$. $E = \frac{1}{2}((2E-1)+1)$ gehört daher zu $R(\mathcal{M}^U)$, und hieraus folgt: $\mathcal{M}^P \subset R(\mathcal{M}^U)$, $R(\mathcal{M}^P) \subset R(\mathcal{M}^U)$, also $R(\mathcal{M}^U) \supset \mathcal{M}$. Damit ist $R(\mathcal{M}^U) = \mathcal{M}$ bewiesen.

3. Sei $\mathcal{M}$ eine beliebige Teilmenge von $\mathcal{B}$. Wir betrachten alle f , für die bei jedem A von $\mathcal{M}$ $Af=0$, $A^*f=0$ gilt — diese bilden offenbar eine abg. lin. M. Der P.O. ihrer Komplementären (vgl. E. Definition 12) heiße E_0 : Zugehörigkeit zu ihr wird also durch $E_0f=0$ gekennzeichnet. Wir definieren:

Definition 4. $\mathcal{M} \subset \mathcal{B}$ sei beliebig. Der soeben konstruierte P.O. E_0 , für den $E_0f=0$ damit gleichwertig ist, daß für alle A von $\mathcal{M}$ $Af=0$, $A^*f=0$ gilt, heiße die Haupteinheit von $\mathcal{M}$. Wir zeigen zuerst:

Satz 3. Für alle A von $\mathcal{M}$, ja von $R(\mathcal{M})$, gilt $E_0A = AE_0 = A$.

Beweis. Gehöre A zu $\mathcal{M}$. Da identisch $E_0(1-E_0)f=0$ ist, ist $A(1-E_0)f=0$ — also $A(1-E_0)=0$, $A=AE_0$. Durch Anwenden auf A^* folgt ebenso $A^*=A^*E_0$, und wenn wir hierauf $*$ anwenden, $A=E_0A$.

Betrachten wir nun alle A mit $E_0A = AE_0 = A$. Sie bilden offenbar einen Ring, und dieser umfaßt, wie wir sahen, $\mathcal{M}$ — also ist er $\supset R(\mathcal{M})$, womit alles bewiesen ist.

Satz 4. E_0 gehört zu $\mathcal{M}'$ und zu $\mathcal{M}''$.

Beweis. Daß alle A von $\mathcal{M}$ mit E_0 vert. sind, sahen wir schon, ferner ist $E_0^*=E_0$ — also gehört E_0 zu $\mathcal{M}'$. Gehöre nun B zu $\mathcal{M}'$. Wenn $E_0f=0$ ist, so ist für alle A von $\mathcal{M}$ $Af=0$, also $BAf=ABf=0$, und ebenso $A^*f=0$, also $BA^*f=A^*Bf=0$ — d. h. $E_0Bf=0$. Da identisch $E_0(1-E_0)f=0$ ist, haben wir also $E_0B(1-E_0)f=0$, d. h. $E_0B(1-E_0)=0$, $E_0B=E_0BE_0$. Wenden wir hierauf $*$ an und ersetzen wir B durch B^* (das ja auch zu $\mathcal{M}'$ gehört), so wird $BE_0=E_0BE_0$ — also $E_0B=BE_0$. Dies gilt für alle B von $\mathcal{M}'$, also gehört E_0 auch zu $\mathcal{M}''$. —

Nun beweisen wir den schon in der Einleitung 3. angekündigten Satz.

Satz 5. $R(\mathcal{M})$ ist die Menge aller A von $\mathcal{M}''$ mit $E_0A = AE_0 = A$.

⁴⁷⁾ Es ist $(2E-1)^*=2E-1$, $(2E-1)^2=4E^2-4E+1=1$.

Beweis. Daß $R(\mathcal{M})$ Teilmenge der genannten Menge ist, wissen wir schon; es ist nur zu zeigen, daß alle derartigen A wirklich zu $R(\mathcal{M})$ gehören.

Seien $\varphi_1, \dots, \varphi_k$ irgendwelche Elemente von $\mathfrak{S}$, die wir fürs Folgende zunächst festhalten, $k = 1, 2, \dots$. Gleichzeitig denken wir uns noch einen Hilbertschen Raum $\bar{\mathfrak{S}}$ gegeben, und wir wählen in $\bar{\mathfrak{S}}$ k unendlichvielde dimensionale abg. lin. M. $\bar{\mathfrak{S}}_1, \dots, \bar{\mathfrak{S}}_k$ aus, die zueinander orth. sind und zusammen die abg. lin. M. $\bar{\mathfrak{S}}$ aufspannen⁴⁸⁾. Die $\bar{\mathfrak{S}}_1, \dots, \bar{\mathfrak{S}}_k$ sind somit selbst Hilbertsche Räume, also mit $\mathfrak{S}$ isomorph⁴⁹⁾ — seien $\Gamma_1, \dots, \Gamma_k$ isomorphe Abbildungen von $\mathfrak{S}$ auf bzw. $\bar{\mathfrak{S}}_1, \dots, \bar{\mathfrak{S}}_k$.

Wir ordnen jetzt jedem A von $\mathcal{B}$ den Punkt $\Gamma_1(A\varphi_1) + \dots + \Gamma_k(A\varphi_k)$ in $\bar{\mathfrak{S}}$ zu⁵⁰⁾, er heiße Bildpunkt von A . Sei $\bar{\mathfrak{L}}$ die Menge der Bildpunkte der A von $R(\mathcal{M})$, offenbar eine lin. M., und $\bar{\mathfrak{M}}$ ihre abg. Hülle (stark in $\bar{\mathfrak{S}}$!), also eine abg. lin. M. Der P. O. von $\bar{\mathfrak{M}}$ heiße $\bar{E}$.

Wenn A ein beschr. O. in $\mathfrak{S}$ ist, so gibt es offenbar einen und nur einen beschr. O. $\bar{A}$ in $\bar{\mathfrak{S}}$, für den

$$\bar{A}(\Gamma_1 f_1 + \dots + \Gamma_k f_k) = \Gamma_1(A f_1) + \dots + \Gamma_k(A f_k)$$

gilt (vgl. Anm.⁵⁰⁾). Es ist offenbar $\bar{A}^* = \bar{A}^*$.

Nun gehöre B zu $R(\mathcal{M})$. Wenn $\bar{f}$ zu $\bar{\mathfrak{L}}$ gehört, so ist $\bar{f} = \Gamma_1(A\varphi_1) + \dots + \Gamma_k(A\varphi_k)$ für ein A von $R(\mathcal{M})$. $\bar{B}f = \Gamma_1(BA\varphi_1) + \dots + \Gamma_k(BA\varphi_k)$, und da BA auch zu $R(\mathcal{M})$ gehört, gehört $\bar{B}f$ auch zu $\bar{\mathfrak{L}}$. $\bar{B}$ bildet also $\bar{\mathfrak{L}}$ auf einen Teil von sich selbst ab — somit erst recht $\bar{\mathfrak{M}}$. Jedes $\bar{E}\bar{f}$ liegt in $\bar{\mathfrak{M}}$, also auch jedes $\bar{B}\bar{E}\bar{f}$, somit gilt identisch $\bar{E}\bar{B}\bar{E}f = \bar{B}\bar{E}f$, d. h. $\bar{E}\bar{B}\bar{E} = \bar{B}\bar{E}$. Wir ersetzen B durch B^* (das ja auch zu $R(\mathcal{M})$ gehört) und wenden $*$ auf diese Gleichung an; wegen $\bar{B}^* = \bar{B}^*$ ist dann $\bar{E}\bar{B}\bar{E} = \bar{E}\bar{B}$; also $\bar{E}\bar{B} = \bar{B}\bar{E}$.

Fassen wir $\bar{E}$ näher ins Auge. Es ist offenbar

$$\bar{E}(\Gamma_1(f_1) + \dots + \Gamma_k(f_k)) = \Gamma_1(g_1) + \dots + \Gamma_k(g_k),$$

wobei die $g_1, \dots, g_k$ lin. und stetig von den $f_1, \dots, f_k$ abhängen. Also ist $g_l = E_{1l}f_1 + \dots + E_{kl}f_k$ ($l = 1, \dots, k$), wobei die E_{jl} ($l, j = 1, \dots, k$)

⁴⁸⁾ Zur Terminologie vgl. etwa E., § I. Die gewünschte Konstruktion erfolgt so: Sei $\bar{\varphi}_1, \bar{\varphi}_2, \dots$ ein vollst. norm. orth. System, wir führen darin eine Doppelindizierung $\bar{\varphi}_{ln}$, $l = 1, \dots, k$, $n = 1, 2, \dots$, ein. Sei $\bar{\mathfrak{S}}_l$ die von $\bar{\varphi}_{l1}, \bar{\varphi}_{l2}, \dots$ aufgespannte abg. lin. M. ($l = 1, \dots, k$) — diese $\bar{\mathfrak{S}}_1, \dots, \bar{\mathfrak{S}}_k$ leisten alles, was wir verlangten.

⁴⁹⁾ Vgl. E., Einleitung III.

⁵⁰⁾ $\Gamma_1 f_1 + \dots + \Gamma_k f_k$ durchläuft offenbar ganz $\bar{\mathfrak{S}}$, und zwar genau einmal, wenn $f_1, \dots, f_k$ unabhängig voneinander $\mathfrak{S}$ durchlaufen.

lauter lin. stetige O. in $\mathfrak{S}$ sind ⁵¹⁾. Daher gilt

$$\begin{aligned}\bar{E} \bar{B} f &= \bar{E} \bar{B} (\Gamma_1(f_1) + \dots + \Gamma_k(f_k)) = \bar{E} (\Gamma_1(B f_1) + \dots + \Gamma_k(B f_k)) \\ &= \Gamma_1(E_{11} B f_1 + \dots + E_{k1} B f_k) + \dots + \Gamma_k(E_{1k} B f_1 + \dots + E_{kk} B f_k), \\ \bar{B} \bar{E} f &= \bar{B} \bar{E} (\Gamma_1(f_1) + \dots + \Gamma_k(f_k)) \\ &= \bar{B} (\Gamma(E_{11} f_1 + \dots + E_{k1} f_k) + \dots + \Gamma_k(E_{1k} f_1 + \dots + E_{kk} f_k)) \\ &= \Gamma_1(B E_{11} f_1 + \dots + B E_{k1} f_k) + \dots + \Gamma_k(B E_{1k} f_1 + \dots + B E_{kk} f_k).\end{aligned}$$

Also muß

$$\begin{aligned}E_{11} B f_1 + \dots + E_{k1} B f_k &= B E_{11} f_1 + \dots + B E_{k1} f_k, \dots, \\ E_{1k} B f_1 + \dots + E_{kk} B f_k &= B E_{1k} f_1 + \dots + B E_{kk} f_k\end{aligned}$$

sein — d. h. allgemein $E_{jl} B = B E_{jl}$. Da B jedes Element von $\mathcal{M}$ sein kann, und dabei noch B, B^* zu $\mathcal{R}(\mathcal{M})$ gehören, also mit E_{jl} vert. sind, gehören alle E_{jl} zu $\mathcal{M}'$.

Die Zugehörigkeit von $\bar{f} = \Gamma_1(f_1) + \dots + \Gamma_k(f_k)$ zu $\bar{\mathcal{M}}$ ist durch $\bar{E} f = f$ charakterisiert, d. h.

$$\begin{aligned}\bar{E} (\Gamma_1(f_1) + \dots + \Gamma_k(f_k)) &= \Gamma_1(f_1) + \dots + \Gamma_k(f_k), \\ \Gamma_1(E_{11} f_1 + \dots + E_{k1} f_k) + \dots + \Gamma_k(E_{1k} f_1 + \dots + E_{kk} f_k) \\ &= \Gamma_1(f_1) + \dots + \Gamma_k(f_k),\end{aligned}$$

$$(E_{11} - 1) f_1 + \dots + E_{k1} f_k = 0, \dots, E_{1k} f_1 + \dots + (E_{kk} - 1) f_k = 0.$$

Und wenn f der Bildpunkt von A ist, $f = \Gamma_1(A \varphi_1) + \dots + \Gamma_k(A \varphi_k)$, so haben wir die Bedingung:

$$\begin{aligned}E_{11} - 1) A \varphi_1 + \dots + E_{k1} A \varphi_k &= 0, \dots, E_{1k} A \varphi_1 + \dots + (E_{kk} - 1) A \varphi_k = 0.\end{aligned}$$

Gehört insbesondere A zu $\mathcal{M}''$, so daß es mit allen E_{jl} (da diese zu $\mathcal{M}'$ gehören) vert. ist, so wird hieraus

$$A[(E_{11} - 1) \varphi_1 + \dots + E_{k1} \varphi_k] = 0, \dots, A[E_{1k} \varphi_1 + \dots + (E_{kk} - 1) \varphi_k] = 0;$$

d. h. wir haben Bedingungen von der Form

$$A \omega_1 = 0, \dots, A \omega_k = 0,$$

mit festen $\omega_1, \dots, \omega_k$.

Für jedes A von $\mathcal{M}$ gehören A, A^* zu $\mathcal{R}(\mathcal{M})$, also ihre Bildpunkte zu $\bar{\mathcal{M}}$, ferner gehören A, A^* auch zu $\mathcal{M}''$ — also ist $A \omega_1 = A^* \omega_1 = \dots = A \omega_k = A^* \omega_k = 0$. Daraus folgt (Definition 4) $E_0 \omega_1 = \dots = E_0 \omega_k = 0$. Wenn also A nur zu $\mathcal{M}''$ gehört, aber dabei $E_0 A = A E_0 = A$ gilt, so haben wir

$$A \omega_1 = A E_0 \omega_1 = 0, \dots, A \omega_k = A E_0 \omega_k = 0,$$

⁵¹⁾ Es ist, wie man leicht ausrechnet, $E_{jl} = \Gamma_j^{-1} P_{\mathfrak{S}_j} \bar{E} \Gamma_l$, und umgekehrt $\bar{E} = \sum_{j,l=1}^k \Gamma_j E_{jl} \Gamma_l^{-1} P_{\mathfrak{S}_l}$.

d. h. der Bildpunkt von A gehört zu $\overline{\mathcal{M}}$. Somit gibt es für jedes $\varepsilon > 0$ ein A' von $\mathbf{R}(\mathcal{M})$, dessen Bildpunkt (das ist der allgemeine Punkt von $\overline{\mathcal{Q}}$) von demjenigen von A eine Entfernung $< \varepsilon$ hat, das hat aber wegen

$$\begin{aligned} & |\Gamma_1(A' \varphi_1) + \dots + \Gamma_k(A' \varphi_k) - \Gamma_1(A \varphi_1) - \dots - \Gamma_k(A \varphi_k)|^2 \\ &= |\Gamma_1((A' - A) \varphi_1) + \dots + \Gamma_k((A' - A) \varphi_k)|^2 \\ &= |\Gamma_1((A' - A) \varphi_1)|^2 + \dots + |\Gamma_k((A' - A) \varphi_k)|^2 \\ &= |(A' - A) \varphi_1|^2 + \dots + |(A' - A) \varphi_k|^2 \end{aligned}$$

die Konsequenz $|(A' - A) \varphi_1| < \varepsilon, \dots, |(A' - A) \varphi_k| < \varepsilon$. D. h.: A' gehört zur starken Umgebung $\mathcal{U}_4(A; \varphi_1, \dots, \varphi_k, \varepsilon)$ von A .

Da nun $\varphi_1, \dots, \varphi_k$ und $\varepsilon > 0$ ganz willkürlich waren, ist A starker Häufungspunkt von $\mathbf{R}(\mathcal{M})$, und da dieses als Ring sogar schwach abg. ist, gehört A dazu.

Damit ist alles bewiesen. —

Aus dem Beweise von Satz 5 läßt sich noch herauslesen:

Zusatz. Wenn eine Menge $\mathcal{M}$ von den Ringeigenschaften $\alpha), \beta)$ der Definition 1 wohl $\alpha)$ besitzt, an Stelle von $\beta)$ aber nur die folgende, weniger weitgehende: wenn A, A^* starke Häufungspunkte von $\mathcal{M}$ sind, so gehört A zu $\mathcal{M}$ (dies ist weniger als die starke Abg., und erst recht als die schwache) — so ist sie dennoch ein Ring.

Beweis. Beim Beweise von Satz 5 wurden von den Eigenschaften von $\mathbf{R}(\mathcal{M})$ nur die folgenden benützt: es ist $> \mathcal{M}, < \mathcal{M}''$, in ihm gilt stets $E_0 A = A E_0 = A$, mit A, B gehören auch $a A, A^*, A + B, AB$ dazu, und unter diesen Prämissen wurde gezeigt, daß jedes A von $\mathcal{M}''$ mit $E_0 A = A E_0 = A$ starker Häufungspunkt von $\mathbf{R}(\mathcal{M})$ ist (die starke Abg. von $\mathbf{R}(\mathcal{M})$, die wir nur brauchten, um hieraus zu folgern, daß A dazu gehört, soll jetzt nicht herangezogen werden). Bei unserem $\mathcal{M}$ können wir daher $\mathbf{R}(\mathcal{M})$ durch $\mathcal{M}$ selbst ersetzen, dann haben wir: jedes A von $\mathcal{M}''$ mit $A E_0 = E_0 A = A$ ist starker Häufungspunkt von $\mathcal{M}$. Da mit A auch A^* so ist, ist A^* auch starker Häufungspunkt von $\mathcal{M}$. Nach der (bisher unbenutzten) zweiten Hälfte der Annahme gehört also A zu $\mathcal{M}$.

Aus Satz 5 folgt daher $\mathbf{R}(\mathcal{M}) \subset \mathcal{M}$, d. h. $\mathbf{R}(\mathcal{M}) = \mathcal{M}$, also ist $\mathcal{M}$ ein Ring, wie behauptet wurde. —

Für Mengen mit der Ringeigenschaft α (Definition 1) ist also starke und schwache Abg. gleichwertig. Dasselbe zeigte in § E. Schmidt für die lin. M.⁵²⁾, unser gegenwärtiges Resultat ist ein Analogon dazu in $\mathcal{B}$.

⁵²⁾ Vgl. Rend. d. Circ. Mat. d. Palermo 25 (1908), S. 57–73.

4. Wir haben soeben $R(\mathcal{M})$ mit Hilfe von E_0 und $\mathcal{M}''$ gekennzeichnet ($\mathcal{M}$ sei wieder eine beliebige Teilmenge von $\mathcal{B}$), nun wollen wir $E_0, \mathcal{M}''$ auf $R(\mathcal{M})$ zurückführen.

Satz 6. E_0 gehört zu $R(\mathcal{M})$, und zwar ist es der größte P.O. in $R(\mathcal{M})$, in dem Sinne, daß jeder P.O. aus $R(\mathcal{M})$ Teil von E_0 ist.

Beweis. Da E_0 zu $\mathcal{M}''$ gehört und $E_0 E_0 = E_0$ ist, gehört es zu $R(\mathcal{M})$ (Satz 5). Wenn E ein P.O. aus $R(\mathcal{M})$ ist, so ist $E E_0 = E$, d. h. E Teil von E_0 (E., Satz 17).

Satz 7. $\mathcal{M}''$ ist die Menge aller $A \vdash a \cdot 1$, oder auch die aller $A + a \cdot (1 - E_0)$ (A aus $R(\mathcal{M})$, a komplex).

Beweis. Da E_0 zu $\mathcal{M}''$ gehört, bedeutet beides dasselbe; wegen $R(\mathcal{M}) \subset \mathcal{M}''$, 1 aus $\mathcal{M}''$, ist es klar, daß die genannte Menge $\subset \mathcal{M}''$ ist — es bleibt zu zeigen, daß jedes B von $\mathcal{M}''$ die Form $A + a \cdot (1 - E_0)$, A aus $R(\mathcal{M})$, hat.

Mit E_0, B gehört auch $E_0 B = B E_0 = A$ zu $\mathcal{M}''$ (da E_0 zu $\mathcal{M}'$, B zu $\mathcal{M}''$ gehört, sind sie vert.), dabei ist $E_0 A = E_0^2 B = E_0 B = A$, $A E_0 = B E_0^2 = B E_0 = A$. Nach Satz 5 gehört also A zu $R(\mathcal{M})$. $C = B - A$ gehört also auch zu $\mathcal{M}''$, und es ist $E_0 C = E_0 B - E_0 A = A - A = 0$, $C E_0 = B E_0 - A E_0 = A - A = 0$. Wenn wir hieraus $C = a \cdot (1 - E_0)$ schließen können, sind wir am Ziele.

Sei F ein P.O., der Teil von $1 - E_0$ ist. Dann ist F zu E_0 orth. (vgl. E., Satz 17), d. h. $E_0 F = F E_0 = 0$. Für jedes A von $R(\mathcal{M})$ ist somit $FA = F \cdot E_0 A = 0$, $AF = A E_0 \cdot F = 0$, d. h. A, F vert. — also für jedes A von $\mathcal{M}$ A, A^* mit F vert., d. h. F gehört zu $\mathcal{M}'$. Somit sind F, C vert. Zusammen mit $E_0 C = C E_0 = 0$ folgt daraus das gewünschte $C = a(1 - E_0)$.⁵³⁾ —

Zum Schluß zeigen wir noch:

Satz 8. Die 1 enthaltenden Ringe, die $\mathcal{M}$ mit $\mathcal{M} = \mathcal{M}''$, und die Mengen $\mathcal{M}'$ sind untereinander identisch.

⁵³⁾ Dies zeigt man so: Sei $E_0 f = 0, f \neq 0$. Wie man sofort einsieht, gehört zu f nach $Fg = \frac{1}{|f|^2} (g, f) \cdot f$ ein P.O., und zwar folgt aus $E_0 f = 0$ $E_0 F = 0$, d. h. E_0, F sind orth. — also F, C vert. Somit ist $Cf = CFf = FCF = \frac{1}{|f|^2} (Cf, f) \cdot f = a_f \cdot f$. Nun wollen wir zeigen, daß die Zahl a_f für alle f dieselbe ist. $a_f = a_{cf}$ ($c \neq 0$) ist klar, d. h. $a_f = a_g$ für lin. abhängige f, g . Sind diese aber unabhängig, so haben wir $C(f+g) = a_{f+g} \cdot (f+g)$, $C(f+g) = Cf + Cg = a_f \cdot f + a_g \cdot g$, $(a_f - a_{f+g}) \cdot f + (a_g - a_{f+g}) \cdot g = 0$, also $a_f = a_g = a_{f+g}$. Also ist $a_f = a$ wirklich konstant, $Cf = af$ für $E_0 f = 0$ ($f \neq 0$ ist offenbar unwesentlich).

Da identisch $E_0(1 - E_0)f = 0$ ist, gilt für alle f $Cf = Cf - CE_0 f = C(1 - E_0)f = a(1 - E_0)f$, also ist $C = a(1 - E_0)$.

Beweis. Wenn ein Ring $\mathcal{M}$ die 1 enthält, so ist für seine Haupteinheit E_0 $E_0 = E_0 1 = 1$, und da für jedes A dann $E_0 A = 1 A = A$, $A E_0 = A 1 = A$ gilt, $R(\mathcal{M}) = \mathcal{M}''$ (Satz 5), d. h. $\mathcal{M} = \mathcal{M}''$. Wenn ein $\mathcal{M}$ der Bedingung $\mathcal{M} = \mathcal{M}''$ genügt, so ist $\mathcal{M} = \mathcal{N}'$ mit $\mathcal{N} = \mathcal{M}'$. Jede Menge $\mathcal{N}'$ ist Ring und enthält die 1 (vgl. Anfang von 1.).

Damit ist die Gleichwertigkeit unserer drei Kriterien erwiesen.

III. Die Abelschen Ringe.

1. Daß unter Voraussetzung der Ringeigenschaft α) (Definition 4) schwache und starke Abg. dasselbe ist, sahen wir schon. Wir betrachten nun Abelsche Mengen, und wollen zeigen, daß bei diesen die Abg.-Eigenschaften noch weiter abgeschwächt werden dürfen.

Sei $\mathcal{M}$ eine zunächst beliebige Menge $\subset \mathcal{B}$, indem wir auf ihre Elemente die Operationen $\alpha A, A^*, A + B, AB$ (beliebig, aber endlich, oft) anwenden, entsteht die Menge $r(\mathcal{M})$ — die offenbar $\subset R(\mathcal{M})$ und $\supset \mathcal{M}$ ist. Wenn wir zu $r(\mathcal{M})$ alle A hinzufügen, für die A und A^* starke Häufungspunkte von $r(\mathcal{M})$ sind, entsteht so eine Menge $r_1(\mathcal{M})$, die $\subset R(\mathcal{M})$ und $\supset \mathcal{M}$ ist, infolge der Symmetrie der Definition mit A auch A^* enthält, und mit A, B auch $\alpha A, A + B, AB$ enthält (weil $r(\mathcal{M})$ es tut, und diese Operationen im Sinne der starken Topologie stetig sind — AB zwar nur in jeder Variablen für sich, fürs Vorliegende genügt aber dies, wie man leicht einsieht). Wenn B und B^* starke Häufungspunkte von $r_1(\mathcal{M})$ sind, so sind sie auch solche von $r(\mathcal{M})$ (denn $r_1(\mathcal{M})$ besteht aus Häufungspunkten von $r(\mathcal{M})$) — daher gehören sie zu $r_1(\mathcal{M})$.

Der Zusatz zu Satz 5 ist also anwendbar: $r_1(\mathcal{M})$ ist ein Ring, woraus $r_1(\mathcal{M}) = R(\mathcal{M})$ folgt. Durch Hinzufügen aller starken bzw. aller schwachen Häufungspunkte entstehe aus $r(\mathcal{M})$ bzw. $r_2(\mathcal{M}), r_3(\mathcal{M})$. Dann ist $r_1(\mathcal{M}) \subset r_2(\mathcal{M}) \subset r_3(\mathcal{M}) \subset R(\mathcal{M})$ klar, also $r_1(\mathcal{M}) = r_2(\mathcal{M}) = r_3(\mathcal{M}) = R(\mathcal{M})$.

Wenn nun $\mathcal{M}$ Abelsch ist (d. h. alle Elemente von $\mathcal{M}$ sowie deren * miteinander vert. sind), so genügt eine noch geringere Erweiterung von $r(\mathcal{M})$, um $R(\mathcal{M})$ zu erhalten, nämlich:

Satz 9. Wenn $\mathcal{M}$ Abelsch ist, so entsteht $R(\mathcal{M})$ aus $r(\mathcal{M})$ bereits durch Hinzufügung aller Limites stark doppelkonvergenter Teilfolgen desselben⁵⁴⁾.

Beweis. Nennen wir diese Menge $\bar{r}(\mathcal{M})$, daß sie $\subset R(\mathcal{M})$ und $\supset \mathcal{M}$ ist, ist klar, ebenso daß sie mit A, B auch $\alpha A, A^*, A + B, AB$

⁵⁴⁾ Dies ist schon darum eine wesentliche Verschärfung, weil weder in der starken noch in der schwachen Topologie das erste Abzählbarkeitsaxiom gilt — d. h. jeder Häufungspunkt auch Limes ist (vgl. § I. 2.). — Ob dieser Satz auch unabhängig vom Abelschen Charakter von $\mathcal{M}$ gilt, konnte nicht entschieden werden.

enthält (weil $r(\mathcal{M})$ so ist, und diese Operationen im Sinne der starken Doppelkonvergenz stetig sind). Wenn also $\bar{r}(\mathcal{M})$ stark abg. ist, so ist es ein Ring (Zusatz zu Satz 5), und daher $= R(\mathcal{M})$. Nur die starke Abg. ist also zu beweisen⁵⁵⁾.

Wenn A starker Häufungspunkt von $\bar{r}(\mathcal{M})$ ist, so ist es einer von $R(\mathcal{M})$, gehört daher zu $R(\mathcal{M})$ — und ist somit normal (§ I, 1.); bei beliebigem $\varphi_1, \dots, \varphi_k, \varepsilon > 0$ gibt es ein A' aus $\bar{r}(\mathcal{M})$ in $\mathcal{U}_3(A; \varphi_1, \dots, \varphi_k, \varepsilon)$, d. h. mit

$$|(A' - A)\varphi_1| < \varepsilon, \dots, |(A' - A)\varphi_k| < \varepsilon,$$

da A' zu $\bar{r}(\mathcal{M})$ gehört, ist es auch normal. Obendrein sind A, A^*, A', A'^* als Elemente von $R(\mathcal{M})$ alle vert. Wir setzen

$$\begin{aligned} \frac{1}{2}(A + A^*) &= B, & \frac{1}{2i}(A - A^*) &= C, & \frac{1}{2}(A' + A'^*) &= B', \\ \frac{1}{2i}(A' - A'^*) &= C', \end{aligned}$$

dies sind dann lauter vert. (beschr.) H. O., und zwar B', C' aus $r(\mathcal{M})$, ferner ist $A' - A = (B' - B) + i(C' - C)$.

Wenn nun B_1, C_1 zwei vert. (beschr.) H. O. sind, so ist

$$\begin{aligned} |(B_1 + iC_1)f|^2 &= |B_1f|^2 + (B_1f, iC_1f) + (iC_1f, B_1f) + |iC_1f|^2 \\ &= |B_1f|^2 + (-iC_1B_1f, f) + (iB_1C_1f, f) + |C_1f|^2 = |B_1f|^2 + |C_1f|^2, \end{aligned}$$

woraus auf Grund der früheren Ungleichheiten

$$|(B' - B)\varphi_1| < \varepsilon, \dots, |(B' - B)\varphi_k| < \varepsilon$$

und

$$|(C' - C)\varphi_1| < \varepsilon, \dots, |(C' - C)\varphi_k| < \varepsilon$$

folgt. Wir bezeichnen $\text{Max}(|B|, |C|)$ mit c und $\text{Max}(|B'|, |C'|)$ mit c' , d. h. es ist stets $|Bf| \leq c \cdot |f|$, $|Cf| \leq c \cdot |f|$, $|B'f| \leq c' \cdot |f|$, $|C'f| \leq c' \cdot |f|$. Während c fest ist, hängt c' von $\varphi_1, \dots, \varphi_k, \varepsilon$ ab — indem wir beide eventuell vergrößern, nehmen wir $c' \geq c > 0$ an.

Jetzt soll das Verhältnis von B, B' näher untersucht werden. Sei $p(x)$ ein Polynom mit den folgenden Eigenschaften:

$$\left. \begin{aligned} -c &< p(x) < \begin{cases} c \\ -c + \delta \end{cases} & \text{für } -c' \leq x \leq -c, \\ x - \delta &< p(x) < \begin{cases} x + \delta \\ c \end{cases} & \text{für } -c \leq x \leq c, \\ -c &< p(x) < c & \text{für } c \leq x \leq c', \end{aligned} \right\}$$

⁵⁵⁾ Man beachte: es ist nicht einmal die Abg. von $\bar{r}(\mathcal{M})$ im Sinne der starken Doppelkonvergenz selbstverständlich.

wobei wir über $\delta > 0$ noch verfügen werden. Man sieht: für $|x| \leq c$ ist $|p(x) - x| < \delta$, für $|x| \leq c'$ ist $|p(x)| < c$, für $|\tau|, |y| \leq c'$ ist $|p(x) - p(y)| < |x - y| + 2\delta$ also $(p(x) - p(y))^2 - (x - y)^2 < 8c'\delta + \delta^2$. Wir wählen nun δ so, daß in der ersten Ungleichheit $|p(x) - x| < \varepsilon$ und in der letzten $(p(x) - p(y))^2 - (x - y)^2 < \varepsilon^2$ gilt ($p(x)$ hängt von c', δ, ε , also mittelbar und unmittelbar von $\varphi_1, \dots, \varphi_k, \varepsilon$ ab).

Aus einem Satze von F. Riesz (vgl. E., Anhang II, Satz 3*, 4*) folgt für alle f $|p(B)f - Bf| \leq \varepsilon \cdot |f|$, ebenso $|p(B')f| \leq c \cdot |f|$. Ferner ist das Polynom $\varepsilon^2 + (x - y)^2 - (p(x) - p(y))^2 = q(x, y)$ im Quadrat $|x|, |y| \leq c'$ stets ≥ 0 , hieraus und aus der Vert. von B, B' sowie $|B|, |B'| \leq c'$, kann man vermöge eines Analogons des soeben benutzten Satzes — auf welches wir gleich zurückkommen werden — folgern, daß der H. O. $q(B, B')$ definit ist. D. h.:

$$(q(B, B')f, f) \geq 0, \quad ((p(B') - p(B))^2 f, f) \leq ((B' - B)^2 f, f) + \varepsilon^2 (f, f),$$

$$|(p(B') - p(B))f|^2 \leq |(B' - B)f|^2 + \varepsilon^2 |f|^2.$$

Für $f = \varphi_1, \dots, \varphi_k$ haben wir also

$$|(p(B') - p(B))\varphi_l| \leq \varepsilon \sqrt{1 + |\varphi_l|^2}.$$

Zusammen mit der früheren Ungleichheit ergibt das

$$|(p(B') - B)\varphi_l| \leq \varepsilon (|\varphi_l| + \sqrt{1 + |\varphi_l|^2}).$$

Wenn wir also $\varepsilon = \eta : \text{Max}_{l=1, \dots, k} (|\varphi_l| + \sqrt{1 + |\varphi_l|^2})$ setzen, $\eta > 0$ beliebig, so liegt $p(B')$ in $\mathcal{U}_s(B; \varphi_1, \dots, \varphi_k, \eta)$. Dabei ist, wie wir sahen, $|p(B')| \leq c$ — d. h. B ist starker Häufungspunkt des in der Kugel $|B''| \leq c$ liegenden Teiles von $\mathcal{r}(\mathcal{M})$, und zwar der H. O. darin. Dann ist es aber (§ I, 3.) auch Limes einer stark konvergenten Folge dortselbst, welche — da es sich ausschließlich um H. O. handelt — sogar doppelkonvergent ist. D. h.: B gehört zu $\overline{\mathcal{r}}(\mathcal{M})$.

Genau so zeigt man, daß auch C dazu gehört, also auch $A = B + iC$, womit alles bewiesen ist.

Es bleibt noch übrig, den vorhin übergangenen Beweis der Verallgemeinerung des F. Rieszschen Satzes nachzuholen. Seien also B, B' zwei vert. H. O. mit $|B|, |B'| \leq c'$, nach E., Anhang II, Satz 11* (man setze $A = B + iB'$) gibt es dann zwei Z. d. E. $E(\lambda), F(\lambda)$, so daß stets

$$(Bf, g) = \int_{-c'}^{c'} \lambda d(E(\lambda)f, g), \quad (B'f, g) = \int_{-c'}^{c'} \lambda d(F(\lambda)f, g)$$

ist, und alle $E(\lambda), F(\mu)$ vert. sind. Wie dort setzen wir

$$\Delta(\lambda, \mu) = E(\lambda) F(\mu) = F(\mu) E(\lambda)$$

und merken uns, daß

$$\Delta(\lambda', \mu') \Delta(\lambda'', \mu'') = \Delta(\text{Min}(\lambda', \lambda''), \text{Min}(\mu', \mu''))$$

ist. Aus den obigen Formeln folgt (alle $\int \int$ sind über $\int_{-c'}^{c'} \int_{-c'}^{c'}$ zu erstrecken):

$$(Bf, g) = \int \int \lambda d(\Delta(\lambda, \mu)f, g), \quad (B'f, g) = \int \int \mu d(\Delta(\lambda, \mu)f, g).$$

Genau wie beim Beweise von Satz 1 (vgl. dazu das Zitat von Anm. 44)) gilt dann

$$(q(B, B')f, g) = \int \int q(\lambda, \mu) d(\Delta(\lambda, \mu)f, g),$$

$$(q(B, B')f, f) = \int \int q(\lambda, \mu) d(\Delta(\lambda, \mu)f, f) = \int \int q(\lambda, \mu) d|\Delta(\lambda, \mu)f|^2.$$

Nun ist im Integrationsgebiet $q(\lambda, \mu) \geq 0$, und es ist $\int \int_{\mathfrak{R}} d|\Delta(\lambda, \mu)f|^2 \geq 0$ für jedes Rechteck $\mathfrak{R}$ (vgl. E., Anhang II, § 5), also ist das letzte Integral rechts ≥ 0 . Da demnach für alle f $(q(B, B')f, f) \geq 0$ ist, ist $q(B, B')$ definit, wie behauptet wurde⁵⁶⁾.

2. Ein Ring $\mathcal{M} = R(A)$ mit einer Erzeugenden A ist dann und nur dann Abelsch, wenn es die Menge (A) ist, d. h. wenn A, A^* vert. sind: also für normales A . Wir wollen nun die schon in Einleitung 3 angekündigte Umkehrung beweisen: jeder Abelsche Ring (d. h. jeder, der nur normale Elemente hat) kann auf die Form $R(A)$ gebracht werden, wo das normale A sogar H. gewählt werden kann.

Satz 10. Wenn A alle normalen O. durchläuft, so durchläuft $R(A)$ alle Abelschen Ringe, und nur diese. Und zwar kann man bei gegebenem $\mathcal{M} = R(A)$ A immer als H. O. wählen.

Beweis. Nach dem vorhin Gesagten genügt es, zu zeigen: zu jedem Abelschen Ring $\mathcal{M}$ existiert ein H. O. A mit $\mathcal{M} = R(A)$.

Zunächst ist $\mathcal{M} = R(\mathcal{M}^P)$ (Satz 2). Nach § I, 4. gibt es eine in $\mathcal{M}^P$ (stark) überall dichte abzählbare Menge $\mathcal{N} < \mathcal{M}^P$. Somit ist $\mathcal{M}^P < R(\mathcal{N})$, $\mathcal{M} = R(\mathcal{M}^P) < R(\mathcal{N})$, und andererseits $R(\mathcal{N}) < R(\mathcal{M}^P) = \mathcal{M}$ — d. h. $\mathcal{M} = R(\mathcal{N})$. Dabei ist $\mathcal{N}$ eine Folge von P. O. $E_1, E_2, \dots$ aus $\mathcal{M}$ — also sind diese vert.

Im folgenden schreiben wir für E Teil von F kürzer $E \leq F$, zwei P. O. E, F sind vergleichbar, wenn $E \leq F$ oder $F \leq E$ ist. Wir haben $\mathcal{M} = R(E_1, E_2, \dots)$ wo alle E_n vert. sind — sie sollen nun (unter Erhaltung der obigen Eigenschaft) so abgeändert werden, daß sie sogar ver-

⁵⁶⁾ Man beachte, daß zu diesem Beweise die Z. d. E. gegebener H. O. herangezogen werden mußten, während umgekehrt aus dem F. Rieszschen Satze deren Existenz bewiesen werden kann. — Über die hier in E. Anhang II verwendeten Stieltjes-Radonschen Doppelintegrale vgl. Wiener Akad. 122, 2 (1913), S. 1295–1438, insbesondere §§ I–II.

gleichbar sind. Genauer: wir werden eine Folge $F_1, F_2, \dots$ von P.O. mit den folgenden Eigenschaften angeben. $F_1 = E_1$; $F_2 \leq F_1 \leq F_3$ und F_1, F_2, F_3 entstehen durch die Operationen $+$, $-$, $\cdot$ aus E_1, E_2 und umgekehrt; $\dots$; $F_{2^{n-1}} \leq F'_1 \leq F_{2^{n-1}+1} \leq F'_2 \leq \dots \leq F'_{2^{n-1}-2} \leq F_{2^{n-1}-1} \leq F'_{2^{n-1}-1} \leq F_{2^n-1}$ (hier sind $F'_1, \dots, F'_{2^{n-1}-1}$ die $F_1, \dots, F_{2^{n-1}-1}$ der durch $\leq$ festgelegten Größe nach geordnet) und $F_1, \dots, F_{2^n-1}$ entstehen durch die Operationen $+$, $-$, $\cdot$ aus $E_1, \dots, E_n$ und umgekehrt; $\dots$. Wenn wir solche F_p haben, so ist es klar, daß jedes F_p durch $+$, $-$, $\cdot$ aus endlich vielen E_n entsteht, und umgekehrt, also ist $R(F_1, F_2, \dots) = R(E_1, E_2, \dots) = \mathcal{M}$; und dabei sind offenbar alle F_p vergleichbar. Wir wollen sie nun herstellen.

Die Konstruktion verläuft induktiv: für $n=1$ ist sie trivial, sie möge also für $n=m-1$ schon geglückt sein, wir führen sie für $n=m$ aus. $F'_1, \dots, F'_{2^{m-1}-1}$ seien wieder die $F_1, \dots, F_{2^{m-1}-1}$ in der $\leq$ -Reihenfolge, wir setzen dann

$$\begin{aligned} F_{2^{m-1}} &= E_m F'_1, & F_{2^{m-1}+1} &= F'_1 + E_m (F'_2 - F'_1), \\ F_{2^{m-1}+2} &= F'_2 + E_m (F'_3 - F'_2), \dots, \\ F_{2^m-2} &= F'_{2^{m-1}-2} + E_m (F'_{2^{m-1}-1} - F'_{2^{m-1}-2}), \\ F_{2^m-1} &= F'_{2^{m-1}-1} + E_m (1 - F'_{2^{m-1}-1}). \end{aligned}$$

Die $F_{2^{m-1}}, \dots, F_{2^m-1}$ sind hier mit Hilfe der $F'_1, \dots, F'_{2^{m-1}-1}, E_m$ ausgedrückt, d. h. mit Hilfe der $F_1, \dots, F_{2^{m-1}-1}, E_m$, also nach Induktionsannahme mit $E_1, \dots, E_{m-1}, E_m$; von den E_n ist nach derselben überhaupt nur E_m zu untersuchen, und da gilt

$$E_m = F_{2^{m-1}} + (F_{2^{m-1}+1} - F'_1) + \dots + (F_{2^m-1} - F'_{2^{m-1}-1}).$$

Schließlich ist auch $F_{2^{m-1}} \leq F'_1 \leq F_{2^{m-1}+1} \leq F'_2 \leq \dots \leq F'_{2^{m-1}-2} \leq F_{2^m-2} \leq F'_{2^{m-1}-1} \leq F_{2^m-1}$ leicht verifizierbar.

Betrachten wir nun die $F_1, F_2, \dots$, ihre Reihenfolge ist:

$$\begin{aligned} F_2 \leq F_1 \leq F_3, & \quad F_4 \leq F_2 \leq F_5 \leq F_1 \leq F_6 \leq F_3 \leq F_7, \\ F_8 \leq F_4 \leq F_9 \leq F_2 \leq F_{10} \leq F_5 \leq F_{11} \leq F_1 \\ & \leq F_{12} \leq F_6 \leq F_{13} \leq F_3 \leq F_{14} \leq F_7 \leq F_{15}, \dots \end{aligned}$$

Dieselbe kann wie folgt gekennzeichnet werden. Wir numerieren die F_p um, indem wir für F_p , $p = 2^{l_1} + 2^{l_2} + \dots + 2^{l_m}$ ($l_1 > l_2 > \dots > l_m \geq 0$, dyadische Entwicklung!), F_ϱ schreiben, $\varrho = \frac{1}{2 \cdot 3^{l_1}} + \frac{2}{3^{l_1-l_2}} + \dots + \frac{2}{3^{l_1-l_m}}$, dann stehen sie in der Größen-Reihenfolge da: $\varrho < \sigma$ zieht $F_\varrho \leq F_\sigma$ nach sich. Dabei durchläuft der Index ϱ , wie man sieht, die Menge aller Mittelpunkte der von der bekannten Cantorschen triadischen Menge ausgelassenen Inter-

valle⁵⁷⁾ — d. h. eine abzählbare Menge, die im Intervalle $0 < x < 1$ liegt, und keinen einzigen ihrer Häufungspunkte enthält.

Wir konstruieren nun eine Z. d. E. $G(\lambda)$ mit den folgenden Eigenschaften: für $\lambda < -1$ $G(\lambda) = 0$, für $\lambda \geq 0$ $G(\lambda) = 1$, für $-1 \leq \lambda < 0$ jedes $G(\lambda)$ (starker) Häufungspunkt der F_ϱ , insbesondere $G(\varrho - 1) = F_\varrho$. Dann ist die Menge $\mathcal{N}$ aller $G(\lambda)$, $\lambda < 0$, und $1 - G(\lambda)$, $\lambda \geq 0$, $\subset \mathcal{M}$, $R(\mathcal{N}) \subset R(\mathcal{M}) = \mathcal{M}$, andererseits umfaßt sie alle F_ϱ , d. h. $F_1, F_2, \dots$, so daß $R(\mathcal{N}) \supset R(F_1, F_2, \dots) = \mathcal{M}$ ist — also $R(\mathcal{N}) = \mathcal{M}$. Wenn also A der beschr. H. O. mit der Z. d. E. $G(\lambda)$ ist (vgl. Anhang I), so ist (Satz 1) $R(A) = R(\mathcal{N}) = \mathcal{M}$, und damit ist alles bewiesen. Daher gilt es, die Z. d. E. $G(\lambda)$ aufzustellen, und zwar genügt es, sie für $-1 \leq \lambda < 0$ zu definieren.

Zunächst sei in jedem von der Cantorschen Menge ausgelassenen (offenen!) Intervalle (vgl. Anm. 57)), wenn dessen Mitte ϱ ist, $\bar{G}(\lambda - 1) = F_\varrho$. Damit ist $\bar{G}(\lambda)$ in einer in $-1 \leq \lambda < 0$ überall dichten offenen Menge definiert — und zwar folgt aus $\lambda < \mu$ $\bar{G}(\lambda) \leq \bar{G}(\mu)$, und aus $\lambda \rightarrow \lambda_0$, $\bar{G}(\lambda)f \rightarrow \bar{G}(\lambda_0)f$ ($\bar{G}(\lambda)$ ist in λ stark halbstetig!). Daher können wir seine Definition auf ganz $-1 \leq \lambda < 0$ so ausdehnen: sei $\lambda_1 > \lambda_2 > \dots$, $\lambda_n \rightarrow \lambda$ eine Folge, für die $\bar{G}(\lambda_1), \bar{G}(\lambda_2), \dots$ schon Sinn haben, wegen $\bar{G}(\lambda_1) \geq \bar{G}(\lambda_2) \geq \dots$ haben die $\bar{G}(\lambda_n)$ einen P. O. G zum starken Limes (vgl. E., Satz 19). $\bar{G}$ hängt nur von λ ab: denn für zwei solche Folgen $\lambda'_1 > \lambda'_2 > \dots$, $\lambda''_1 > \lambda''_2 > \dots$, können wir aus Teilfolgen eine neue $\lambda'_m > \lambda''_n$, $> \lambda''_m > \lambda'_n$ kombinieren, die auch konvergieren muß — also sind die Limes dieselben. Wenn $\bar{G}(\lambda)$ Sinn hat, so ist nach obigem $G = \bar{G}(\lambda)$ — wir bezeichnen allgemein G mit $G(\lambda)$.

Für $\mu < \lambda$ wählen wir $\lambda_1 > \lambda_2 > \dots$, $\lambda_n \rightarrow \lambda$, $\mu_1 > \mu_2 > \dots$, $\mu_n \rightarrow \mu$ mit $\mu_n < \lambda_n$, dann ist $\bar{G}(\mu_n) \leq \bar{G}(\lambda_n)$, also aus Stetigkeitsgründen $G(\mu) \leq G(\lambda)$. Ferner ist bei gegebenem f für geeignetes λ_n $|\bar{G}(\lambda_n)f|^2 \leq |G(\lambda)f| + \varepsilon$, also für $\lambda \leq \lambda' \leq \lambda_n$ ($\lambda_n > \lambda$) erst recht

$$|G(\lambda')f|^2 \leq |G(\lambda_n)f|^2 = |\bar{G}(\lambda_n)f|^2 \leq |G(\lambda)f|^2 + \varepsilon,$$

$$|G(\lambda')f - G(\lambda)f|^2 = |G(\lambda')f|^2 - |G(\lambda)f|^2 \leq \varepsilon$$

(vgl. E., Satz 16) — d. h. für $\lambda' \geq \lambda$, $\lambda' \rightarrow \lambda$ $G(\lambda')f \rightarrow G(\lambda)f$. Also ist $G(\lambda)$ die gewünschte Z. d. E., und wir sind am Ziele. —

⁵⁷⁾ Die genannte Menge umfaßt alle Zahlen $\sum_{s=1}^{\infty} \frac{\alpha_s}{3^s}$, wo alle $\alpha_s = 0, 2$ sind. Sie ist bekanntlich perfekt und nirgends dicht, und läßt die Intervalle

$$\sum_{s=1}^k \frac{\alpha_s}{3^s} + \frac{1}{3^{k+1}} < x < \sum_{s=1}^k \frac{\alpha_s}{3^s} + \frac{2}{3^{k+1}}$$

aus, die Mittelpunkte derselben sind die Punkte $\sum_{s=1}^k \frac{\alpha_s}{3^s} + \frac{1}{2 \cdot 3^k}$ — eine andere Schreibweise für unsere ϱ .

Wenn $\mathcal{M}$ ein Abelscher Ring ist, so ist er nach dem soeben bewiesenen Satz 10 gleich $R(A)$, wobei A ein H.O. ist; und nach Satz 9 ist $R(A)$ die Menge der starken Doppellimites von $r(A)$. $r(A)$ aber ist (wegen $A = A^*$) im vorliegenden Falle die Menge aller $p(A)$, wo $p(x)$ alle Polynome mit $p(0) = 0$ durchläuft. Also: $\mathcal{M}$ ist die Menge der Limites aller stark doppelkonvergenten Folgen $p_1(A), p_2(A), \dots$ ($p_1(x), p_2(x), \dots$ für $x = 0$ verschwindende Polynome). Durch eine allgemeinere Fassung des Begriffes $f(A)$, als wir sie bisher benutzten (d. i. Ausdehnung auf unstetige $f(x)$), könnte man zeigen: $\mathcal{M}$ ist die Menge aller $f(A)$ ($f(x)$ eine beliebige für $x = 0$ verschwindende Funktion, für die die erweiterte Definition gilt). Wir wollen jedoch hier auf diese Dinge nicht eingehen.

IV. Allgemeine Eigenschaften der (unbeschränkten) normalen Operatoren.

1. Wir gehen zu unserem zweiten Gegenstande über, der Betrachtung unbeschränkter O., und der Definition der Normalität für solche. Daher müssen wir zunächst beliebige (nicht notwendig überall sinnvolle oder stetige — wir verlangen vorläufig nicht einmal Lin. und Abg.) O. betrachten, die wir $R, S, \dots$ nennen werden (im Gegensatze zu den $A, B, \dots$ aus $\mathcal{B}$). Wir definieren:

Definition 5. aR ist derjenige O., der für die f Sinn hat, für die Rf sinnvoll ist, und zwar ist sein Wert dann $a \cdot Rf$; $R \pm S$ ist derjenige O., der für die f Sinn hat, für die Rf und Sf Sinn haben, und zwar ist sein Wert dann $Rf \pm Sf$; RS ist derjenige O., der für die f Sinn hat, für die Sf und $R(Sf)$ Sinn haben, und zwar ist sein Wert dann $R(Sf)$ ⁵⁸⁾.

R, A sind vert. (es ist wesentlich, daß A zu $\mathcal{B}$ gehört), wenn RA Fortsetzung von AR ist (d. h. bei sinnvollem Rf auch $R(Af)$ sinnvoll ist — $Af, A(Rf)$ sind es sowieso —, und zwar gleich $A(Rf)$)⁵⁹⁾.

⁵⁸⁾ Daß bei dieser Definition das kommutative und assoziative Gesetz der Addition sowie das assoziative der Multiplikation gelten, ist klar. Vom distributiven gilt immer $(R \pm S) \cdot T = R \cdot T \pm S \cdot T$, dagegen $R \cdot (S \pm T) = R \cdot S \pm R \cdot T$ nur für lin. und überall sinnvolles R (ist R nur lin., so ist die linke Seite Fortsetzung der rechten). Es ist $R + 0 = R$, $1 \cdot R = R \cdot 1 = R$, $R \cdot 0 = 0$ (wenn Rf für $f = 0$ verschwindet) dagegen $0 \cdot R$ nur dort definiert, wo R es ist. Analog gestaltet sich das Verhältnis von $+$ und $-$.

Man sieht: im Unbeschränkten sind die Operationen $+, -, \cdot$ nicht mehr so durchsichtig wie im Beschränkten.

⁵⁹⁾ Nach E., Einleitung V. ist R Fortsetzung von S , wenn Rf überall sinnvoll ist, wo Sf es ist, und dort beide gleich sind. — Wenn auch R überall sinnvoll ist (z. B. zu $\mathcal{B}$ gehört), so bedeutet die jetzige Definition der Vert. $RA = AR$, d. h. die gewöhnliche Vert. Für beliebige R käme es kaum in Frage, $RA = AR$ zu verlangen, da dann R mit 0 nicht vert. wäre (vgl. Anm. ²¹⁾, ⁵⁸⁾). Die Unsymmetrie unserer Definition (zwischen R, A) zeigt, daß eine direkte Definition der Vert., für beliebige R, S , auf Schwierigkeiten stoßen muß.

Nun können wir die Definition 3 auf Mengen $\mathcal{M}$ beliebiger O. erweitern

Definition 3'. Wenn $\mathcal{M}$ eine Menge beliebiger O. ist (nicht mehr notwendig $\subset \mathcal{B}$!), so sei $\mathcal{M}'$ die Menge aller A von $\mathcal{B}$, für die A und A^* mit jedem R von $\mathcal{M}$ vert. sind.

$\mathcal{M}'$ ist also auf alle Fälle $\subset \mathcal{B}$, $\mathcal{M}'', \mathcal{M}''', \dots$ fallen somit stets unter die alte Definition. Was bleibt nun von den Eigenschaften von $\mathcal{M}'$ noch erhalten?

Man sieht sofort: 1 gehört allenfalls zu $\mathcal{M}'$, und mit A, B auch A^*, AB ; wenn ferner alle Elemente von $\mathcal{M}$ lin. sind, so gehören mit A, B auch $aA, A+B$ dazu. Seien nun alle Elemente von $\mathcal{M}$ abg., A ein starker Häufungspunkt von $\mathcal{M}'$ und R ein Element von $\mathcal{M}$. Wenn Rf Sinn hat, so gibt es ein A'_n von $\mathcal{M}'$ in $\mathcal{U}_s \left(A; f, Rf, \frac{1}{n} \right)$, d. h. $A'_n f \rightarrow Af$, $RA'_n f = A'_n Rf \rightarrow ARf$, wegen der Abg. von R ist somit $R(Af)$ sinnvoll und gleich ARf . Somit sind A, R vert.

Wenn also A, A^* starke Häufungspunkte von $\mathcal{M}'$ sind, so gehört A zu $\mathcal{M}'$. Nunmehr nehmen wir zusammenfassend an, daß alle Elemente von $\mathcal{M}$ abg. lin. sind. Dann sind nach dem vorher Gesagten alle Prämissen des Zusatzes zu Satz 5 erfüllt: also $\mathcal{M}'$ ein Ring. Da 1 offenbar dazugehört, ist $\mathcal{M}' = \mathcal{M}'''$ (Satz 8).

Bei beliebigem $\mathcal{M}$ können wir die Resultate vom Anfang von § II, 1 auf $\mathcal{M}' \subset \mathcal{B}$ anwenden: $\mathcal{M}' \subset \mathcal{M}''' = \mathcal{M}^v = \mathcal{M}^{vii} = \dots$, $\mathcal{M}'' = \mathcal{M}^{iv} = \mathcal{M}^{vi} = \dots$, $\mathcal{M}' = \mathcal{M}'''$ dagegen können wir nicht schließen, da die Relation $\mathcal{M} \subset \mathcal{M}''$ fehlt (es ist ja $\mathcal{M}'' \subset \mathcal{B}$, aber nicht notwendig $\mathcal{M} \subset \mathcal{B}$!). Es gilt aber, wie wir sahen, wenn $\mathcal{M}$ aus lauter abg. lin. O. besteht. —

Schließlich konstatieren wir, welche Elemente $\mathcal{M}'^p$ und $\mathcal{M}'^u$ haben. Zu $\mathcal{M}'^p$ gehört jeder P.O. E , der mit allen R von $\mathcal{M}$ vert. ist ($E = E^*$!) — d. h. jeder der alle R von $\mathcal{M}$ reduziert (vgl. E., Definition 13). Zu $\mathcal{M}'^u$ gehört jedes unitäre U , für welches $U, U^* = U^{-1}$ beide mit jedem R von $\mathcal{M}$ vert. sind, d. h. wo RU Forts. von UR und RU^{-1} Forts. von $U^{-1}R$ ist. Man sieht sofort, daß dies bedeutet: $RU = UR$ (auch die Definitionsbereiche identisch!), oder auch: $U^{-1}RU = R$.

2. Wir führen nun den Begriff des konjugierten Paares (kurz: konj. Paar) von O. ein.

Definition 5. Zwei O. R, R^* bilden ein konj. Paar, wenn sie denselben Definitionsbereich haben, dieser die abg. lin. M. $\mathfrak{S}$ aufspannt, und in ihm durchweg

$$(Rf, g) = (f, R^*g), \quad (f, Rg) = (R^*f, g)$$

(vgl. den Anfang der Anm. ¹²)) gilt.

Man sieht sofort: Mit R, R^* bilden auch R^*, R ein konj. Paar. Ein R kann nicht mit zwei verschiedenen R_1^*, R_2^* je ein konj. Paar bilden (für R_1^*, R_2^* sind die Definitionsbereiche gleichlautend vorgeschrieben, ebenso $(R_1^*f, g), (R_2^*f, g)$ für eine die abg. lin. M. $\S$ aufspannende Menge der g : also die R_1^*f, R_2^*f selbst) — ebensowenig zwei verschiedene R_1, R_2 mit demselben R^* (vgl. die erste Bemerkung). Das Zeichen $*$ ist motiviert, wenn man beachtet, daß für ein A aus $\mathcal{B}$ gerade A, A^* ein konj. Paar bilden. R ist dann und nur dann ein H. O., wenn R, R ein konj. Paar bilden.

Ein konj. Paar S, S^* ist Forts. eines anderen R, R^* , wenn der Definitionsbereich von S, S^* den von R, R^* umfaßt, und im letzteren stets $Rf = Sf, R^*f = S^*f$ gilt⁶⁰). Wenn beide nicht identisch sind, sprechen wir (ganz wie bei H. O., vgl. E. § II) von eig. Forts.; ein konj. Paar ohne eig. Forts. ist max. (= maximal).

Wir bemerken noch: jede Forts. eines H. O. ist ein H. O., d. h. wenn R, R^* die Forts. S, S^* hat, so folgt aus $R = R^*S = S^*$. In der Tat: seien Sf, Rg sinnvoll, dann ist

$$(Sf, g) = (f, S^*g) = (f, R^*g) = (f, Rg) = (f, Sg) = (S^*f, g),$$

und da diese g die abg. lin. M. $\S$ aufspannen, $Sf = S^*f$ — also $S = S^*$.

Man erkennt ganz leicht (ganz wie bei H. O., vgl. E. § II, Satz 9), daß jedes konj. Paar zu einem solchen fortgesetzt werden kann, das aus zwei lin. O. besteht — dagegen ist es nicht ohne weiteres zu sehen, ob auch der abg. Charakter derselben so erreichbar ist⁶¹). Wir gehen hierauf nicht näher ein.

3. Ein A von $\mathcal{B}$ ist normal, wenn A, A^* vert. sind, d. h. wenn (A) Abelsch ist, oder, was dasselbe bedeutet (vgl. § II, 1), wenn $(A)''$ Abelsch ist. Dies übertragen wir auf beliebige R (da stets $(R)'' \subset \mathcal{B}$ ist!):

Definition 6. Sei R, R^* ein konj. Paar, es ist normal, wenn $(R)''$ (eine Menge $\subset \mathcal{B}$!) Abelsch ist.

Auf Grund des vorhin Gesagten ist es klar, daß dies für O. aus $\mathcal{B}$ die alte Definition ist.

Satz 11. R, R^* ist dann und nur dann normal, wenn R^*, R es ist.

Beweis: Wir zeigen $(R)' = (R^*)'$, hieraus folgt ja $(R)'' = (R^*)''$, und damit die Behauptung; ferner genügt es aus Symmetriegründen

⁶⁰) Jede der letztgenannten Relationen folgt aus der anderen: denn das auf den Definitionsbereich der R, R^* eingeschränkte S, S^* ist noch immer ein konj. Paar.

⁶¹) Mit den genannten Methoden läßt sich bloß eine Art gemeinsame Abg. beider O. R, R^* des Paares erreichen: aus $f_n \rightarrow f, Rf_n \rightarrow f^*, R^*f_n \rightarrow f^{**}$ folgt die Sinnvollheit von Rf, R^*f , und daß diese bzw. gleich f^*, f^{**} sind.

$(R)' < (R^*)'$ zu beweisen, d. h. daß, wenn A, A^* mit R vert. sind, sie auch mit R^* vert. sein müssen.

Sei also R^*f sinnvoll, dann ist Rf, RAf, RA^*f es auch, also R^*Af, R^*A^*f gleichfalls. Wenn g ebenso ist, so haben wir:

$$(R^*Af, g) = (Af, Rg) = (f, A^*Rg) = (f, RA^*g) = (R^*f, A^*g) = (AR^*f, g), \\ (R^*A^*f, g) = (A^*f, Rg) = (f, ARg) = (f, RA^*g) = (R^*f, Ag) = (A^*R^*f, g),$$

und da diese g die abg. lin. M. $\mathfrak{S}$ aufspannen, $R^*Af = AR^*f, R^*A^*f = A^*R^*f$. Somit sind A, A^* mit R^* vert., wie behauptet wurde.

Satz 12. R, R (R ein H. O.!) ist, wenn R lin. abg. angenommen wird, dann und nur dann normal, wenn R hypermax. ist⁶²⁾.

Beweis: Zunächst stellen wir einige allgemeine Betrachtungen über abg. lin. H. O. an. Wir bilden die Cayleysche Transformierte U von R , $\mathfrak{E}, \mathfrak{F}$ seien ihr Definitionsbereich bzw. Wertvorrat (vgl. E. § V, Satz 2), beide abg. lin. M. Der P. O. von $\mathfrak{E}$ sei E (vgl. E § III, Satz 13).

Sei A mit R vert. Jedes φ mit sinnvollem $U\varphi$ hat die Form $Rf + i f$, daher hat RAf Sinn, und es ist $A\varphi = ARf + i Af = RAf + i Af$. Also ist auch $U(A\varphi)$ sinnvoll, und zwar $= RAf - i Af = ARf - i Af = A(U\varphi)$, d. h. U, A sind vert. Sei nun umgekehrt A mit U vert. Jedes f mit sinnvollem Rf hat die Form $\varphi - U\varphi$, daher hat $UA\varphi$ Sinn, und es ist $Af = A\varphi - AU\varphi = A\varphi - UA\varphi$. Also ist auch $R(Af)$ sinnvoll, und zwar $= i(A\varphi + UA\varphi) = i(A\varphi + AU\varphi) = A(Rf)$, d. h. R, A sind vert. Aus alledem folgt aber $(R)' = (U)'$, und somit erst recht $(R)'' = (U)''$.

Jetzt zeigt sich die Hinreichendheit der Bedingung sofort: wenn R hypermax. ist, so ist U unitär (vgl. E., Satz 35), also aus $\mathcal{B}$, ferner mit $U^* = U^{-1}$ vert. — also ist U, U^* normal, und mit ihm (wegen $(R)'' = (U)''$) auch R, R . Es bleibt übrig, die Notwendigkeit zu prüfen.

Sei also R, R normal. Gehöre A zu $(R)'$ — da es dann auch zu $(U)'$ gehört, ist mit Uf auch UAf sinnvoll, d. h. mit f gehört auch Af zu $\mathfrak{E}$. Ef gehört stets zu $\mathfrak{E}$, also auch AEf , also ist $EAEf = AEf$, — d. h. $EAE = AE$. Durch Anwenden von $*$ und Ersetzen von A durch A^* (das auch zu $(R)'$ gehört) folgt hieraus $EAE = EA$, also ist $AE = EA$, d. h. A, E vert. Somit ist A mit U, E vert., also auch mit UE — dies ist wichtig, weil $V = UE$ überall sinnvoll ist und so zu $\mathcal{B}$ gehört. Da dies für jedes A von $(R)'$ (das mit A auch A^* enthält!) gilt, gehört V zu $(R)''$.

⁶²⁾ Abkürzung für hypermaximal, vgl. E., Definition 9. — Wenn R nicht lin. abg. ist, so betrachte man $\tilde{R}$ (E. § II, Satz 9). Da es H. ist, ist $\tilde{R}, \tilde{R}$ ein konj. Paar; da alles, was mit R vert. ist, es auch mit $\tilde{R}$ ist, ist $(R)' < (\tilde{R})', (R)'' > (\tilde{R})''$, also auch $(\tilde{R})''$ Abelsch. Somit ist auch $\tilde{R}, \tilde{R}$ normal und daher unser Satz auf abg. lin. $\tilde{R}$ anwendbar.

Also auch V^* , und weil $(R)''$ Abelsch ist, sind V, V^* vert. Nun ist für alle f, g

$$(V^* V f, g) = (V f, V g) = (U(E f), U(E g)) = (E f, E g) = (E f, g),$$

also $V^* V = V V^* = E$, und auch dieses muß zu $(R)''$ gehören. Daher sind V, E vert., was die Konsequenz

$$E V = V E = U E \cdot E = U E = V$$

hat. Wenn φ zu $\mathfrak{E}$ gehört, so ist $\varphi = E \varphi$, $U \varphi = U E \varphi = V \varphi = E V \varphi$, d. h. auch $U \varphi$ gehört zu $\mathfrak{E}$, und daher auch $\varphi - U \varphi$. Aber die $\varphi - U \varphi$ müssen überall dicht liegen (E., Satz 24), also ist $\mathfrak{E}$ eine überall dichte abg. lin. M. — d. h. $\mathfrak{E} = \mathfrak{S}$, $E = 1$. Hieraus folgt:

$$V = U \cdot 1 = U, \quad V^* V = V V^* = 1, \quad \text{also} \quad U^* U = U U^* = 1,$$

d. h. U unitär, und daher R hypermax.

Wir weisen noch einmal auf die beim Beweise dieses Satzes nebenbei gefundene Tatsache hin, daß für einen hypermax. H. O. R und seine (unitäre) Cayleysche Transformierte $(R)' = (U)'$ gilt — und daher $(R)'' = (U)''$, $(R)''' = (U)'''$, $(U)''$ ist Abelsch, also $(U)'' < (U)''' = (U)'$, $(R)'' < (R)'$, und dabei $(R)' = (R)''' = \dots$, $(R)'' = (R)^{IV} = \dots$ (weil es für U so ist, oder nach § IV, 1).

Satz 13. R sei (wie vorhin) ein hypermax. H. O., U seine Cayleysche Transformierte, $E(\lambda)$ seine Z. d. E., und $\mathcal{E}$ die Menge aller $E(\lambda)$. Dann ist $(R)' = (U)' = \mathcal{E}'$ — und daher

$$(R)'' = (U)'' = \mathcal{E}'', \quad (R)''' = (U)''' = \mathcal{E}''', \dots$$

Beweis: Aus der ersten Gleichung folgt alles, und da $(R)' = (U)'$ schon feststeht, ist nur $(U)' = \mathcal{E}'$ zu beweisen — was jedenfalls aus $R(U) = R(\mathcal{E})$ folgt⁶³). Es genügt also zu zeigen: U gehört zu $R(\mathcal{E})$, alle $E(x)$ ⁶⁴ gehören zu $R(U)$.

Daß U zu $R(\mathcal{E})$ gehört, beweist man auf Grund der Relation

$$(U f, g) = \int_0^1 e^{2\pi i x} d(E(x) f, g)$$

wörtlich, wie das Entsprechende in der ersten Hälfte des Beweises von Satz 1 für den dortigen beschr. H. O. A gezeigt wird. Um einzusehen, daß umgekehrt $E(x)$ zu $R(U)$ gehört, vergegenwärtigen wir uns, wie diese $E(x)$ in E., Anhang II konstruiert wurden. Die $E(x)$ waren Differenzen

⁶³) Allgemein gilt: $\mathcal{M} < R(\mathcal{M}) < \mathcal{M}''$, $\mathcal{M}' > (R(\mathcal{M}))' > \mathcal{M}''' = \mathcal{M}'$, $(R(\mathcal{M}))' = \mathcal{M}'$, d. h. $R(\mathcal{M})$ bestimmt $\mathcal{M}'$.

⁶⁴) Wir schreiben $E(x)$, $0 < x < 1$, statt $E(\lambda)$, $-\infty < \lambda < \infty$, wobei der Zusammenhang $\lambda = -\operatorname{ctg} \pi x$, $x = -\frac{1}{\pi} \operatorname{arc} \operatorname{ctg} \lambda$ (vgl. E., Beweis von Satz 36) besteht.

starker Limites gewisser $F(\varrho)$ ($0 < \varrho < 1$, vgl. § 6 dortselbst), es genügt also, wenn diese zu $R(U)$ gehören; die $F(\varrho)$ aber entstanden durch Addieren, Subtrahieren, Multiplizieren aus den $E(U, a, b)$, ($-\infty < a, b < \infty$, vgl. §§ 4 u. 5 dortselbst), also sind nur diese von Interesse. $E(U, a, b)$ aber war Produkt von P. O. aus den Z. d. E. von $\frac{U+U^*}{2}$ und $\frac{U-U^*}{2i}$ ⁶⁵⁾, und diese gehören nach unserem Satz 2 $R(U)$ an, wenn $\frac{U+U^*}{2}, \frac{U-U^*}{2i}$ ihm angehören. Dieses ist klar: $R(U)$ umfaßt U, U^* , also auch die letztgenannten O.

4. Nach diesen Vorbereitungen nehmen wir unseren eigentlichen Gegenstand wieder in Angriff. R, R^* sei wieder ein konj. Paar.

Satz 14. R, R^* sei normal. Wir bilden die zwei H. O. $S_1 = \frac{R+R^*}{2}$, $S_2 = \frac{R-R^*}{2i}$, und dann $\tilde{S}_1, \tilde{S}_2$ (vgl. Anfang Anm. ⁶³⁾), die letzteren sind hypermax. Die O. $\tilde{R} = \tilde{S}_1 + i\tilde{S}_2$, $\tilde{R}^* = \tilde{S}_1 - i\tilde{S}_2$ (beide definiert im Durchschnitt der Definitionsbereiche von $\tilde{S}_1, \tilde{S}_2$) bilden ein konj. Paar, und zwar ein normales. $\tilde{R}, \tilde{R}^*$ sind Forts. von R, R^* .

Beweis: Wir bilden S_1, S_2 wie vorgeschrieben, der H. Charakter ist evident, und dann $\tilde{S}_1, \tilde{S}_2$. Gehöre A zu $\mathcal{B}$. Wenn es mit R, R^* vert. ist, so ist es auch mit S_1, S_2 vert., also auch mit $\tilde{S}_1, \tilde{S}_2$ — hieraus folgert man $(R, R^*)' < (\tilde{S}_1)'$ und $(\tilde{S}_2)'$. Wegen $(R)' = (R^*)'$ (vgl. den Beweis von Satz 11) ist $(R)' = (R, R^*)'$, also $(R)' < (\tilde{S}_1)'$ und $(\tilde{S}_2)'$, $(R)'' > (\tilde{S}_1)''$ und $(\tilde{S}_2)''$. Mit $(R)''$ sind somit $(\tilde{S}_1)''$, $(\tilde{S}_2)''$ Abelsch, d. h. die konj. Paare $\tilde{S}_1, \tilde{S}_1$ und $\tilde{S}_2, \tilde{S}_2$ normal — also nach Satz 12 $\tilde{S}_1, \tilde{S}_2$ hypermax.

Da $\tilde{S}_1, \tilde{S}_2$ bzw. Forts. von S_1, S_2 sind, ist $\tilde{R}, \tilde{R}^*$ eine solche von R, R^* , daß $\tilde{R}, \tilde{R}^*$ ein konj. Paar ist, ist klar, da $\tilde{S}_1, \tilde{S}_2$ H. O. sind. Die Normalität erkennen wir so: Es ist $(\tilde{R})' = (\tilde{R}^*)'$, also beides $= (\tilde{R}, \tilde{R}^*)'$, und dieses offenbar $> (\tilde{S}_1, \tilde{S}_2)'$, welches seinerseits, wie wir schon wissen, $> (R)'$ ist. Also: $(\tilde{R})' > (R)'$, $(\tilde{R})'' < (R)''$, mit $(R)''$ ist daher auch $(\tilde{R})''$ Abelsch, also $\tilde{R}, \tilde{R}^*$ normal.

Satz 15. $\tilde{R}f$ (und $\tilde{R}^*f$) ist dann und nur dann sinnvoll, wenn zwei f^*, f^{**} existieren, so daß für alle g mit sinnvollem Rg (und R^*g) $(f, Rg) = (f^{**}, g)$, $(f, R^*g) = (f^*, g)$ gilt. Und zwar ist dann $Rf = f^*$, $R^*f = f^{**}$.

⁶⁵⁾ $E(U, a, b)$ (a. a. O. heißt $U A$) ist als $P_{\mathfrak{R}(R, a)} \cdot P_{\mathfrak{R}(S, b)}$ ($R = \frac{U+U^*}{2}, S = \frac{U-U^*}{2i}$) definiert, und es kommt nur auf die in Satz 8*, 9* dortselbst angegebenen Eigenschaften dieser P. O. an, welche den bzw. Z. d. E. offenbar zukommen.

Beweis: Die Notwendigkeit ist klar:

$$(f, Rg) = (f, \tilde{R}g) = (\tilde{R}^*f, g), \quad (f, R^*g) = (f, \tilde{R}^*g) = (\tilde{R}f, g).$$

Die Hinreichendheit erkennen wir so: $\frac{f^* + f^{**}}{2}$ und $\frac{f^* - f^{**}}{2i}$ sind nach Definition dem f durch S_1 bzw. S_2 zugeordnet (vgl. E., Definition 8, 9), also auch durch $\tilde{S}_1, \tilde{S}_2$. Da die letzteren hypermax. sind, ist $\tilde{S}_1 f, \tilde{S}_2 f$ sinnvoll, und zwar gleich $\frac{f^* + f^{**}}{2}, \frac{f^* - f^{**}}{2i}$. Also sind auch $\tilde{R}f, \tilde{R}^*f$ sinnvoll, und zwar gleich f^*, f^{**} .

Satz 16. Jede Forts. T, T^* von R, R^* (konj. Paar!, normal braucht sie nicht zu sein) hat $\tilde{R}, \tilde{R}^*$ zur Forts. $\tilde{R}, \tilde{R}^*$ ist (selbst als bloßes conj. Paar) max., und zwar die einzige max. Forts. von R, R^* .

Beweis: Die zwei letzten Behauptungen folgen aus der ersten, betrachten wir daher nur diese. Seien Tf, T^*f sinnvoll, wenn wir sie gleich f^*, f^{**} setzen, so sind die Prämissen von Satz 15 erfüllt: bei sinnvollem Rg, R^*g ist

$$(f, Rg) = (f, Tg) = (T^*f, g), \quad (f, R^*g) = (f, T^*g) = (Tf, g).$$

Also sind auch $\tilde{R}f, \tilde{R}^*f$ sinnvoll, und zwar gleich Tf, T^*f . D. h. $\tilde{R}, \tilde{R}^*$ ist Forts. von T, T^* . —

Bei normalen conj. Paaren von O. herrschen also viel einfachere Verhältnisse als bei den H. O. (vgl. E.): so ist die max. Forts. hier stets auf eine und nur eine Weise möglich. Wir werden uns daher im folgenden stets auf max. Paare R, R^* beschränken, diese sind nach Satz 16 durch $\tilde{R} = R, \tilde{R}^* = R^*$ charakterisiert. Für diese werden wir das Analogon der Hilbertschen Spektraldarstellung aufstellen.

V. Spektralform der normalen Operatoren.

1. Wir untersuchen, wie angekündigt, normale max. conj. Paare R, R^* .

Satz 17. Sei R, R^* ein normales max. conj. Paar. Es gibt dann eine Schar von P.O. $\Delta(z)$ (z durchläuft alle komplexen Zahlen) mit den folgenden Eigenschaften:

$$\alpha) \Delta(z') \cdot \Delta(z'') = \Delta(z'') \cdot \Delta(z') = \Delta(z)$$

$$(\Re z = \min(\Re z', \Re z''), \Im z = \min(\Im z', \Im z'')).$$

$\beta) \Delta(z)f \rightarrow \Delta(z_0)f$ für alle f und $z \rightarrow z_0$, wenn dabei $\Re z \geq \Re z_0$, $\Im z \geq \Im z_0$ bleibt, und zwar ist die Konvergenz sowohl bei festem $f, \Re z$ und beliebigem $\Im z$, als auch bei festem $f, \Im z$ und beliebigem $\Re z$, gleichmäßig.

$\gamma)$ $\Delta(z)f \rightarrow 0$ bzw. $\rightarrow f$ für alle f , wenn $\Re z \rightarrow -\infty$ oder $\Im z \rightarrow -\infty$ bzw. wenn $\Re z \rightarrow +\infty$ und $\Im z \rightarrow +\infty$.

(Auf Grund dieser Eigenschaften α) bis γ) heiße eine Schar von P. O. eine komplexe Z. d. E.)

$\delta)$ Rf (und R^*f) hat dann und nur dann Sinn, wenn

$$\iint |z|^2 d|\Delta(z)f|^2$$

endlich ist. ($\iint$ ist über die ganze komplexe Zahlenebene zu erstrecken. Da $\iint_{\Re} d|\Delta(z)f|^2$ für jedes Rechteck $\Re \geq 0$ ist⁶⁶) und der Integrand $|z|^2$ stets ≥ 0 ist, ist das obige Integral seiner Natur nach entweder konvergent und endlich, oder eigentlich divergent und $+\infty$. Diese letztere Möglichkeit soll ausgeschlossen werden.)

$\varepsilon)$ Im genannten Falle gilt für alle g

$$(Rf, g) = \iint z d(\Delta(z)f, g), \quad (R^*f, g) = \iint \bar{z} d(\Delta(z)f, g).$$

(Die Integrale sind für solche f — und beliebige g — absolut konvergent⁶⁷.)

(Auf Grund von δ) bis ε) sollen R, R^* und $\Delta(z)$ zusammengehörig heißen.)

⁶⁶) Wenn das Rechteck $\Re$ die Ecken $a' + b'i, a' + b''i, a'' + b''i, a'' + b'i$ ($a' \leq a'', b' \leq b''$) hat, so ist

$$\begin{aligned} \iint_{\Re} |\Delta(z)f|^2 &= |\Delta(a' + b'i)f|^2 - |\Delta(a' + b''i)f|^2 \\ &\quad + |\Delta(a'' + b''i)f|^2 - |\Delta(a'' + b'i)f|^2 \end{aligned}$$

und wegen $|Ef|^2 = (f, Ef)$ (E ein P. O.)

$$= |\{\Delta(a' + b'i) - \Delta(a' + b''i) + \Delta(a'' + b''i) - \Delta(a'' + b'i)\}f|^2 \geq 0,$$

falls der O. $\{\dots\}$ ein P. O. ist. In der Tat ist nach $\alpha)$

$$\begin{aligned} \Delta(a' + b'i) - \Delta(a' + b''i) + \Delta(a'' + b''i) - \Delta(a'' + b'i) \\ = \Delta(a'' + b''i)(1 - \Delta(a' + b''i))(1 - \Delta(a'' + b'i)), \end{aligned}$$

womit auch dies in Evidenz gesetzt ist.

⁶⁷) Alle diese Ausführungen stehen in weitestgehender Analogie zu denjenigen von E., Satz 36, mit denen man sie vergleichen möge. Der Vollständigkeit halber sei hier ein Beweis der absoluten Konvergenz von $\iint z d(\Delta(z)f, g)$ ($\iint \bar{z} d(\Delta(z)f, g)$ erledigt sich ebenso) gegeben.

Seien $\Re_1, \dots, \Re_k$ irgendwelche Rechtecke (ohne gemeinsame innere Punkte), die zusammen einen endlichen Teil der Ebene ausfüllen; $\zeta_1, \dots, \zeta_k$ bzw. Punkte derselben. Dann gilt ($h = 1, \dots, k$, E_h sei der nach Anm. ⁶⁶) fürs Rechteck $\Re_h$ berechnete P. O.):

$$\begin{aligned} \left| \iint_{\Re_h} d(\Delta(z)f, g) \right| &= |(E_h f, g)| = |(E_h f, E_h g)| \leq |E_h f| \cdot |E_h g| \\ &= \sqrt{(E_h f, f) \cdot (E_h g, g)} = \sqrt{\iint_{\Re_h} d|\Delta(z)f|^2 \cdot \iint_{\Re_h} d|\Delta(z)g|^2}, \end{aligned}$$

(Fortsetzung der Fußnote ⁶⁷) auf nächster Seite.)

ζ) Wenn die H. O. $\tilde{S}_1, \tilde{S}_2$ nach Satz 14 gebildet werden, so ist $\tilde{S}_1 f$ bzw. $\tilde{S}_2 f$ dann und nur dann sinnvoll, wenn

$$\iint (\Re z)^2 d|\Delta(z)f|^2 \quad \text{bzw.} \quad \iint (\Im z)^2 d|\Delta(z)f|^2$$

endlich ist. (Über die Natur dieser Integrale vgl. S. 411 und Anm. ⁶⁶)).

η) Im genannten Falle gilt für alle g

$$(\tilde{S}_1 f, g) = \iint \Re z \cdot d(\Delta(z)f, g), \quad (\tilde{S}_2 f, g) = \iint \Im z \cdot d(\Delta(z)f, g)$$

(Die absolute Konvergenz der Integrale ist gesichert, vgl. ε) und Anm. ⁶⁷)).

Beweis. Da $\tilde{S}_1, \tilde{S}_2$ hypermax. sind, gehören zu ihnen zwei Z. d. E. $E(\lambda), F(\lambda)$ (vgl. E., Satz 36, sowie unsere Anm. ⁴³)). Nach Definition ist $\tilde{S}_1 f$ bzw. $\tilde{S}_2 f$ dann sinnvoll, wenn $\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2$ bzw. $\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2$ endlich ist, und zwar ist dann für jedes g

$$(\tilde{S}_1 f, g) = \int_{-\infty}^{\infty} \lambda d(E(\lambda)f, g) \quad \text{bzw.} \quad (\tilde{S}_2 f, g) = \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g).$$

Die Menge der $E(\lambda)$ bzw. der $F(\lambda)$ heiße $\mathcal{E}$ bzw. $\mathcal{F}$. Jedes $E(\lambda)$ gehört zu $\mathcal{E}$, also zu $\mathcal{E}''$, also nach Satz 13 zu $(\tilde{S}_1)''$. Schon beim Beweise von Satz 14 sahen wir aber, daß $(\tilde{S}_1)'' \subset (R)''$ ist, also ist $\mathcal{E} \subset (R)''$. Ebenso ergibt sich $\mathcal{F} \subset (R)''$. Da $(R)''$ Abelsch ist, sind die Elemente von $\mathcal{E}$ mit denen von $\mathcal{F}$ vert.: d. h. jedes $E(a)$ mit jedem $F(b)$. Durch $\Delta(a + ib) = E(a)F(b) = F(b)E(a)$ definieren wir also eine Schar von P. O. $\Delta(z)$. Man überzeugt sich leicht davon, daß a) bis γ) gelten, d. h. daß dies eine komplexe Z. d. E. ist.

Ferner gehen die Integrale

$$\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2, \quad \int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2, \quad \int_{-\infty}^{\infty} \lambda d(E(\lambda)f, g), \quad \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g)$$

und daraus folgt (vgl. die Abschätzungen in E., Beweis von Satz 36, sowie unsere Anm. ⁴⁵))

$$\begin{aligned} \left| \sum_{h=1}^k \zeta_h \iint_{\Re h} d(\Delta(z)f, g) \right| &\leq \sum_{h=1}^k |\zeta_h| \sqrt{\iint_{\Re h} d|\Delta(z)f|^2 \cdot \iint_{\Re h} d|\Delta(z)f|^2} \\ &\leq \sqrt{\sum_{h=1}^k |\zeta_h|^2 \iint_{\Re h} d|\Delta(z)f|^2} \cdot \sqrt{\sum_{h=1}^k \iint_{\Re h} d|\Delta(z)g|^2}. \end{aligned}$$

Da aber die Integrale

$$\iint |z|^2 d|\Delta(z)f|^2 \quad \text{und} \quad \iint d|\Delta(z)g|^2 = |g|^2$$

endlich sind, also absolut konvergieren (vgl. die Bemerkung zu δ)), folgt hieraus dasselbe für $\iint z d(\Delta(z)f, g)$.

in

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} a^2 d|\Delta(a+ib)f|^2, \quad \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} b^2 d|\Delta(a+ib)f|^2,$$

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} a d(\Delta(a+ib)f, g), \quad \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} b d(\Delta(a+ib)f, g)$$

über, also wenn $z = a + ib$ gesetzt wird, in diejenigen aus ζ) bis η). Also sind auch diese erfüllt.

Wegen $R = \bar{R}$, $R^* = \bar{R}^*$ hat Rf (und R^*f) dann und nur dann Sinn, wenn $\tilde{S}_1 f$, $\tilde{S}_2 f$ Sinn haben — d. i. wenn beide Integrale aus ζ) endlich sind. Nach dem, was dort über ihre wesentlich nicht-negative Natur gesagt wurde, bedeutet dies, daß ihre Summe endlich ist; und wegen $|z|^2 = (\Re z)^2 + (\Im z)^2$ ist das genau die Aussage von δ). ε) folgt ohne weiteres aus η).

Satz 18. Die Schar $\Delta(z)$ ist bereits durch die Bedingungen α) bis ε) eindeutig festgelegt: d. h. zu jedem max. konj. Paar R, R^* gehört eine und nur eine komplexe Z. d. E. $\Delta(z)$.

Beweis. Sei $\Delta(z)$ die beim Beweise von Satz 17 konstruierte Schar und $\Delta'(z)$ eine andere, die α) bis ε) auch erfüllt — es gilt $\Delta(z) = \Delta'(z)$ für alle z zu beweisen.

Zunächst folgt η) aus δ) zumindest für diejenigen f , für die Rf Sinn hat, d. h. $\tilde{S}_1 f$, $\tilde{S}_2 f$ beide Sinn haben — für diese f gilt es also auch mit $\Delta'(z)$. Das auf den Definitionsbereich von R eingeschränkte $\tilde{S}_1$ bzw. $\tilde{S}_2$ ist aber S_1 bzw. S_2 (beide haben denselben Definitionsbereich): also handelt es sich von den f , für die $S_1 f$, $S_2 f$ Sinn haben.

Da $\Delta'(a+ib)$ bei wachsendem b (und festem a !) im Sinne der P. O.-Relation $\leq$ nie fällt (E. § III, Satz 14), so existiert für $b \rightarrow \infty$ ein starker Limes (E. § III, Satz 19): $\Delta'(a+ib) \rightarrow E'(a)$. Ebenso existiert bei festem b und $a \rightarrow \infty$ ein starker Limes: $\Delta'(a+ib) \rightarrow F'(b)$. Die $E'(a)$, $F'(b)$ sind P. O., die (ebenso wie $\Delta(a+ib)$, aus Stetigkeitsgründen) bei wachsendem a, b nie fallen; dabei ist für $a' \geq a, b' \geq b$ nach α)

$$\Delta'(a' + ib) \Delta'(a + ib') = \Delta'(a + ib') \Delta'(a' + ib) = \Delta'(a + ib),$$

also durch Grenzübergang ($a', b' \rightarrow +\infty$)

$$F'(b) E'(a) = E'(a) F'(b) = \Delta'(a + ib).$$

Aus

$$\Delta'(a + ib) f \rightarrow \begin{cases} E'(a) f & \text{für } b \rightarrow +\infty \\ 0 & \text{für } b \rightarrow -\infty \end{cases} \quad \Delta'(a + ib) f \rightarrow \begin{cases} F(b) f & \text{für } a \rightarrow +\infty \\ 0 & \text{für } a \rightarrow -\infty \end{cases}$$

folgt man sofort

$$\iint \Re z \cdot d(\Delta'(z)f, g) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} a d(E'(a)F'(b)f, g) = \int_{-\infty}^{\infty} a d(E'(a)f, g),$$

$$\iint \Im z \cdot d(\Delta'(z)f, g) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} b d(E'(a)F'(b)f, g) = \int_{-\infty}^{\infty} b d(F'(b)f, g),$$

in dem Sinne, daß jede rechte Seite konvergiert, wenn es die entsprechende linke tut (die beiden Gleichungen gelten unabhängig voneinander).

Weiter schließt man mühelos aus α) bis γ), daß $E'(\lambda)$, $F'(\lambda)$ Z. d. E. sind, es mögen zu ihnen die H. O. T_1 bzw. T_2 gehören (vgl. E., Satz 36). Nach dem bisher Gesagten ist, wenn $T_1 f$ und $S_1 g$ beide Sinn haben,

$$(f, S_1 g) = \int_{-\infty}^{\infty} \lambda d(f, E(\lambda)g) = \int_{-\infty}^{\infty} \lambda d(E(\lambda)f, g) = (T_1 f, g),$$

d. h. f ist Erweiterungselement von S_1 , und zwar wird ihm durch S_1 $T_1 f$ zugeordnet — also auch durch $\tilde{S}_1$. Aber $\tilde{S}_1$ ist hypermax.: also ist $\tilde{S}_1 f$ sinnvoll und gleich $T_1 f$. D. h.: $\tilde{S}_1$ ist Forts. von T_1 und, da auch T_1 hypermax. ist, $\tilde{S}_1 = T_1$. Ebenso zeigt man $\tilde{S}_2 = T_2$. Da nun $\tilde{S}_1, \tilde{S}_2$ die bzw. Z. d. E. $E(\lambda), F(\lambda)$ haben, muß $E'(\lambda) = E(\lambda), F'(\lambda) = F(\lambda)$ sein. Hieraus folgt aber

$$\Delta'(a + ib) = E'(a)F'(b) = E(a)F(b) = \Delta(a + ib),$$

d. h. $\Delta'(z) = \Delta(z)$, was zu beweisen war.

Satz 19. Zu jeder komplexen Z. d. E., d. h. jeder α) bis γ) erfüllenden Schar $\Delta(z)$ gehört ein und nur ein konj. Paar R, R^* (nach δ) bis ε), und dieses ist max. und normal.

Beweis. Wenn es ein solches Paar gibt, gibt es nur eines: denn δ) legt seinen Definitionsbereich fest, und ε) die $(Rf, g), (R^*f, g)$, d. h. die Rf, R^*f selbst. Bezüglich der Existenz ist noch zu sagen: wenn wir irgendein O.-Paar R, R^* nach δ), ε) haben (und der Definitionsbereich überall dicht ist), so bilden diese nach δ), ε) von selbst ein konj. Paar — nur die Max. und Normalität ist dann noch zu prüfen.

Die Existenz der R, R^* nach δ), ε) steht fest, sobald wir dieses wissen: wenn $\iint |z|^2 d|\Delta(z)f|^2$ endlich ist, so gibt es zwei f^*, f^{**} , so daß für alle g

$$\iint z d(\Delta(z)f, g) = (f^*, g), \quad \iint \bar{z} d(\Delta(z)f, g) = (f^{**}, g)$$

gilt (dieses ist dann Rf bzw. R^*f). Wir nennen die Integrale auf den linken Seiten vorübergehend $L^*(g), L^{**}(g)$ (f sei fest gedacht). Da offenbar

$$\begin{aligned} L^*(a_1 g_1 + \dots + a_n g_n) &= \bar{a}_1 L^*(g_1) + \dots + \bar{a}_n L^*(g_n), \\ L^{**}(a_1 g_1 + \dots + a_n g_n) &= \bar{a}_1 L^{**}(g_1) + \dots + \bar{a}_n L^{**}(g_n) \end{aligned}$$

gilt, sind wir nach dem Satz von F. Riesz (vgl. E., Anm. ⁵²) am Ziele, wenn wir eine Konstante c mit $|L^*(g)| \leq c \cdot |g|$, $|L^{**}(g)| \leq c \cdot |g|$ finden. Aber aus der Abschätzung von Anm. ⁶⁷) folgt

$$|L^*(g)| \leq \sqrt{\iint |z|^2 |d\Delta(z)f|^2} \sqrt{\iint d|\Delta(z)g|^2} = \sqrt{\iint |z|^2 d|\Delta(z)f|^2} \cdot |g|,$$

und ebenso

$$|L^{**}(g)| \leq \sqrt{\iint |z|^2 d|\Delta(z)f|^2} \cdot |g|,$$

womit alles Gewünschte erreicht ist.

Der Definitionsbereich, d. h. die f mit endlichem $\iint |z|^2 d|\Delta(z)f|^2$, ist überall dicht. Dies zeigt man so: Wie beim Beweise von Satz 18 zu $\Delta'(z)$ die $E'(\lambda)$, $F'(\lambda)$ gebildet wurden, bilden wir jetzt zu $\Delta(z)$ $E(\lambda)$, $F(\lambda)$ zwei Z. d. E. $\iint (z)^2 d|\Delta(z)f|^2$ ist endlich, wenn $\iint (\Re z)^2 d|\Delta(z)f|^2$, $\iint (\Im z)^2 d|\Delta(z)f|^2$ es sind, d. h. wenn $\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2$, $\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2$ (vgl. dort) es sind. Das trifft gewiß für alle f mit $E(\lambda)f$ und

$F(\lambda)f = \begin{cases} f & \text{für } \lambda \geq c \\ 0 & \text{für } \lambda \leq -c \end{cases} \quad (c \text{ beliebig, aber fest}) \text{ zu}^{68)} \text{ — also für jedes } f = (E(c) - E(-c))(F(c) - F(-c))g. \text{ Da dieses für } c \rightarrow \infty \text{ gegen } g \text{ konvergiert und } g \text{ beliebig ist, ist die betrachtete Menge wirklich überall dicht.}$

Es bleibt noch zu zeigen, daß R , R^* max. und normal ist. Zunächst die Normalität. Sei $\mathcal{D}$ die Menge aller $\Delta(z)$, da diese H. und vert. sind, ist $\mathcal{D}$ Abelsch, also auch $\mathcal{D}''$. Daher genügt es, $(R)'' \subset \mathcal{D}''$ zu beweisen, also erst recht $(R)' \supset \mathcal{D}'$. Dies ist erwiesen, wenn wir zeigen: wenn A mit allen $\Delta(z)$ vert. ist, so ist es auch mit R vert.

Sei also Rf sinnvoll, wir behaupten: $R(Af)$ ist sinnvoll und gleich $A(Rf)$. Wir müssen also erstens aus der Endlichkeit von $\iint |z|^2 d|\Delta(z)f|^2$ auf diejenige von $\iint |z|^2 d|\Delta(z)Af|^2$ schließen und dann etwa $(R(Af), g) = (A(Rf), g)$ (für alle g !), d. h. $(Af, R^*g) = (Rf, A^*g)$, also $\iint z d(Af, \Delta(z)g) = \iint z d(\Delta(z)f, A^*g)$ beweisen.

A ist beschränkt, d. h. stets $|Af| \leq c \cdot |f|$, hieraus folgt für jedes Rechteck $\Re$ (wenn $E_{\Re}$ der nach Anm. ⁶⁶) dazu konstruierte P.O. ist)

$$\iint_{\Re} d|\Delta(z)Af|^2 = |E_{\Re}Af|^2 = |AE_{\Re}f|^2 \leq c^2 |E_{\Re}f|^2 = c^2 \iint_{\Re} d|\Delta(z)f|^2;$$

wegen $\lambda^2 \geq 0$ ist also $\iint \lambda^2 d|\Delta(z)Af|^2 \leq c^2 \iint \lambda^2 d|\Delta(z)f|^2$, womit die erste Behauptung bewiesen ist. Die zweite dagegen ist nach

$$(Af, \Delta(z)g) = (\Delta(z)Af, g) = (A\Delta(z)f, g) = (\Delta(z)f, A^*g)$$

trivial.

⁶⁸) Denn dann können die Integrale $\int_{-\infty}^{\infty}$ durch $\int_{-c}^c$ ersetzt werden, und dann ist der Integrand λ^2 beschränkt.

Betrachten wir nun die Max. Diese besagt $\tilde{R} = R$, $\tilde{R}^* = R^*$, d. h.: da der Durchschnitt der Definitionsbereiche von $\tilde{S}_1, \tilde{S}_2$ (vgl. Satz 14) der Definitionsbereich von $\tilde{R}, \tilde{R}^*$ ist, ist zu verlangen, daß R (und damit R^*) in ihm überall sinnvoll ist. Wir ziehen wieder die Z. d. E. $E(\lambda), F(\lambda)$ von vorhin heran, zu ihnen mögen die (hypermax.) H. O. T_1, T_2 gehören. Offenbar ist T_1 Forts. von S_1 , also (da es abg. lin. ist) auch von $\tilde{S}_1$ und, da dieses hypermax. ist, $T_1 = S_1$ — ebenso zeigt man $T_2 = S_2$. Daher sind im Durchschnitt der Definitionsbereiche von $\tilde{S}_1 f, \tilde{S}_2 f$ beide $T_1 f, T_2 f$ sinnvoll, d. h. die Integrale

$$\begin{aligned} \int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)f|^2 &= \iint (\Re z)^2 d|\Delta(z)f|^2, \\ \int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2 &= \iint |\Im z|^2 d|\Delta(z)f|^2 \end{aligned}$$

endlich — also auch ihre Summe $\iint |z|^2 d|\Delta(z)f|^2$. Somit ist Rf sinnvoll, also alles bewiesen.

Die Sätze 17 bis 19 zeigen, daß die max. normalen konj. Paare von O. den komplexen Z. d. E. in ebenso ein-eindeutiger Weise entsprechen, wie die hypermax. H. O. den reellen Z. d. E. (nach E., Satz 36). Wegen der allgemeinen Eindeutigkeit der max. Forts. (und dem Fehlen eines Analogons der Hypermax.) sind die Verhältnisse sogar etwas einfacher.

2. Wir stellen noch einige einfache Eigenschaften der max. normalen konj. Paare fest.

Satz 20. R, R^* sei ein max. normales konj. Paar, $\Delta(z)$ seine komplexe Z. d. E., $\mathcal{D}$ die Menge aller $\Delta(z)$. Dann ist $(R)' = \mathcal{D}', (R)'' = \mathcal{D}'', (R)''' = \mathcal{D}''', \dots$

Beweis. Es genügt offenbar, die erste Gleichung zu beweisen, und $(R)' > \mathcal{D}'$ sahen wir schon beim vorigen Beweise — es bleibt also noch $(R)' < \mathcal{D}'$ zu zeigen. Schon beim Beweise von Satz 14 kam $(R)' < (\tilde{S}_1)'$ und $(\tilde{S}_2)'$, d. h. $(R)' < (\tilde{S}_1, \tilde{S}_2)'$ vor, es genügt also, $(\tilde{S}_1, \tilde{S}_2)' < \mathcal{D}'$ einzusehen. Die Z. d. E. von $\tilde{S}_1$ bzw. $\tilde{S}_2$ sei $E(\lambda)$ bzw. $F(\lambda)$, die Menge aller $E(\lambda)$ bzw. $F(\lambda)$ $\mathcal{E}$ bzw. $\mathcal{F}$; nach Satz 13 ist $(\tilde{S}_1)' = \mathcal{E}', (\tilde{S}_2)' = \mathcal{F}'$, also $(\tilde{S}_1, \tilde{S}_2)' = (\mathcal{E}, \mathcal{F})'$.

Wenn nun A (aus $\mathcal{B}$) mit allen Elementen von $\mathcal{E}$ und $\mathcal{F}$ vert. ist, so ist es auch mit jedem $\Delta(a + ib) = E(a)F(b)$ vert., d. h. mit jedem Element von $\mathcal{D}$. Somit ist $(\mathcal{E}, \mathcal{F})' < \mathcal{D}'$. Damit ist alles bewiesen.

Satz 21. R, R^* sei ein normales konj. Paar, dann gilt allgemein $|Rf| = |R^*f|$. Wenn R, R^* auch max. ist, so ist sowohl R als auch R^* ein abg. lin. O.

Beweis. Da R, R^* allenfalls eine max. Forts. besitzt (Satz 16), genügt es, die erste Behauptung für max. R, R^* zu beweisen. Sei also $\Delta(z)$

seine komplexe Z. d. E. Wenn Rf (und R^*f) Sinn hat, so gilt

$$\begin{aligned}(Rf, \Delta(a + ib)g) &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a' + ib') d(\Delta(a' + ib')f, \Delta(a + ib)g) \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a' + ib') d(\Delta(a + ib)\Delta(a' + ib')f, g) \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a' + ib') d(\Delta(\text{Min}(a, a') + i\text{Min}(b, b'))f, g) \\ &= \int_{-\infty}^a \int_{-\infty}^b (a' + ib') d(\Delta(a + ib)f, g),\end{aligned}$$

also insbesondere

$$\begin{aligned}(\Delta(a + ib)f, Rf) &= \int_{-\infty}^a \int_{-\infty}^b (a' - ib') d(f, \Delta(a + ib)f) \\ &= \int_{-\infty}^a \int_{-\infty}^b (a' - ib') d|\Delta(a + ib)f|^2.\end{aligned}$$

Daraus folgt weiter

$$\begin{aligned}|Rf|^2 &= (Rf, Rf) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a + ib) d(\Delta(a + ib)f, Rf) \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a + ib) d\left[\int_{-\infty}^a \int_{-\infty}^b (a' - ib') d|\Delta(a' + ib')f|^2\right]^{69)} \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a + ib)(a - ib) d|\Delta(a + ib)f|^2 \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a^2 + b^2) d|\Delta(a + ib)f|^2.\end{aligned}$$

Ebenso zeigt man

$$|R^*f|^2 = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a^2 + b^2) d|\Delta(a + ib)f|^2.$$

Damit ist $|Rf| = |R^*f|$ in der Tat bewiesen.

Gehen wir nun zur zweiten Behauptung über. Daß R wie R^* lin. ist, folgt z. B. aus Satz 15 (wegen der Max. ist ja $R = \tilde{R}$, $R^* = \tilde{R}^*$). Wir zeigen jetzt die Abg. von R — die von R^* wird ebenso bewiesen. Sei $f_n \rightarrow f$, alle Rf_n (und daher R^*f_n) sinnvoll, $Rf_n \rightarrow f^*$. Also erfüllen die Rf_n die Cauchysche Konvergenzbedingung ($|R(f_m - f_n)| \rightarrow 0$ für $m, n \rightarrow \infty$), die R^*f_n daher auch (weil $|R(f_m - f_n)| = |R^*(f_m - f_n)|$ ist) — also existiert ein f^{**} mit $R^*f_n \rightarrow f^{**}$. Hieraus folgt sofort, daß f, f^*, f^{**} die Prämissen von Satz 15 erfüllen, d. h. daß Rf sinnvoll und gleich f^* ist. Also ist R abg., wie behauptet wurde.

⁶⁹⁾ Zu dieser Integral-Umformung vgl. E., Anm. ⁶⁵⁾.

Anhang I.

Der Vollständigkeit halber soll hier die bekannte (Hilbertsche) Kennzeichnung der Z. d. E. der beschr. H. O. gegeben werden. Sei also R hypermax., $E(\lambda)$ eine Z. d. E.

Nach E. § II, Satz 12 γ ist für die Beschr. von R $|(Rf, f)| \leq c \cdot |f|^2$ (c fest!) charakteristisch, und dabei ist

$$(Rf, f) = \int_{-\infty}^{\infty} \lambda d(E(\lambda)f, f) = \int_{-\infty}^{\infty} \lambda d|E(\lambda)f|^2.$$

Wenn $E(\lambda) = \begin{cases} 1 & \text{für } \lambda \geq c \\ 0 & \text{für } \lambda \leq -c \end{cases}$ ist, so können wir $\int_{-\infty}^{\infty}$ durch $\int_{-c}^c$ ersetzen; dann bleibt der Integrand absolut $\leq c$, also das Integral ($|E(\lambda)f|^2$ nimmt nie ab!) absolut $\leq c \int_{-\infty}^{\infty} d|E(\lambda)f|^2 = c \cdot |f|^2$ — d. h. R ist beschr.

Ist umgekehrt $E(\lambda) \neq 0$ bzw. $\neq 1$ für beliebig kleine bzw. große λ , also für absolut beliebig große λ nicht konstant (wegen $E(\lambda)f \rightarrow 0$ bzw. f für $\lambda \rightarrow -\infty$ bzw. $+\infty$ kommt ja kein anderer konstanter Wert für $E(\lambda)$ als 0 bzw. 1 in Frage!), so kommt $E(\lambda) \neq E(\mu)$ und $\lambda < \mu < -D$ oder $D < \lambda < \mu$ für jedes D vor. Somit ist $E(\mu) - E(\lambda)$ ein P. O. und $\neq 0$; sei $f \neq 0$ ein Punkt seiner abg. lin. M., d. h. $(E(\mu) - E(\lambda))f = f$.

Dann ist $E(\lambda')f = \begin{cases} f & \text{für } \lambda' \geq \mu \\ 0 & \text{für } \lambda' \leq \lambda \end{cases}$.⁷⁰⁾ Also haben wir

$$\begin{aligned} |(Rf, f)| &= \left| \int_{-\infty}^{\infty} \lambda' d|E(\lambda')f|^2 \right| = \left| \int_{\lambda}^{\mu} \lambda' d|E(\lambda')f|^2 \right| \geq D \int_{\lambda}^{\mu} d|E(\lambda')f|^2 \\ &= D \int_{-\infty}^{\infty} d|E(\lambda)f|^2 = D \cdot |f|^2. \end{aligned}$$

Für $D > c$ ist dies $> c \cdot |f|^2$, d. h. R unbeschr.

Es ist ein leichtes, analog die Beschr. für max. normale konj. Paare R, R^* an der komplexen Z. d. E. $A(z)$ abzulesen. Sie ist durch $|Rf| \leq c \cdot |f|$ (c fest!) definiert, man legt vorteilhaft die folgende Relation zugrunde (nach Satz 21):

$$|Rf|^2 = \iint |z|^2 d|A(z)f|^2.$$

⁷⁰⁾ Es ist ja

$$|E(\mu)f|^2 - |E(\lambda)f|^2 = |(E(\mu) - E(\lambda))f|^2 = |f|^2,$$

$$|E(\mu)f|^2 \leq |f|^2, \quad |E(\lambda)f|^2 \geq 0$$

(vgl. E. § III), also $|E(\mu)f| = |f|$, $|E(\lambda)f| = 0$. Erst recht ist daher für $\lambda' \geq \mu$ bzw. $\leq \lambda$ $|E(\lambda')f| = |f|$ bzw. 0 — d. h. $E(\lambda')f = f$ bzw. 0 (vgl. a. a. O.), wie behauptet.

Anhang II.

1. Es soll eine Anwendung der Resultate der §§ IV—V gemacht werden, insbesondere auf Laurentsche Matrizen und deren Verallgemeinerungen.

Sei $\mathcal{P}$ eine Teilmenge von $\mathcal{B}$, für welche $\mathcal{P}'$ Abelsch ist. Wenn R, R^* ein konj. Paar ist, so sagen wir, daß es von der Klasse $\mathcal{P}$ ist, falls $(R)' \supset \mathcal{P}$ gilt — d. h. wenn R mit allen A, A^* vert. ist, falls A ganz $\mathcal{P}$ durchläuft.

Da dann $(R)'' \subset \mathcal{P}'$ ist, ist $(R)''$ Abelsch, d. h. R, R^* normal; nach Satz 16 hat es somit eine einzige max. Forts., das ebenfalls normale $\bar{R}, \bar{R}^*$. Beim Beweise von Satz 14 sahen wir $(\bar{R})' \supset (R)'$, also $\supset \mathcal{P}$ — d. h. auch $\bar{R}$ ist von der Klasse $\mathcal{P}$. Somit genügt es, max. normale R, R^* zu betrachten. Als solches werde R, R^* vorausgesetzt, seine komplexe Z. d. E. sei $\Delta(z)$, $\mathcal{D}$ die Menge aller $\Delta(z)$. Wegen $(R)' = \mathcal{D}'$ (Satz 20) ist R, R^* dann und nur dann von der Klasse $\mathcal{P}$, wenn $\mathcal{D}' \supset \mathcal{P}$ ist. (Wegen $(\Delta(z)) \subset \mathcal{D}$, $(\Delta(z))' \supset \mathcal{D}' \supset \mathcal{P}$ gehört also dann jedes $\Delta(z)$ zur Klasse $\mathcal{P}$ — da $\mathcal{D}'$ der Durchschnitt aller $(\Delta(z))'$ ist, ist dies auch hinreichend.) Aus $\mathcal{D}' \supset \mathcal{P}$ folgt $\mathcal{P}' \supset \mathcal{D}'' \supset \mathcal{D}$, aus $\mathcal{P}' \supset \mathcal{D}$ $\mathcal{D}' \supset \mathcal{P}'' \supset \mathcal{P}$, also ist auch $\mathcal{D} \subset \mathcal{P}'$, d. h. die Zugehörigkeit aller $\Delta(z)$ zu $\mathcal{P}'$, notwendig und hinreichend. —

Wir stellen nun ein $\mathcal{P}$ von der oben genannten Art auf die folgende Weise her. Sei $\mathcal{G}$ eine abzählbar-unendliche Abelsche Gruppe (die Methode ist auch auf kontinuierliche Gruppen $\mathcal{G}$ verallgemeinerbar, was hier nicht erörtert werde), $\alpha, \beta, \dots$ ihre Elemente. Sei $\varphi_\alpha, \varphi_\beta, \dots$ ein vollst. norm. orth. System (in $\mathfrak{H}$), das wir nicht mit den Zahlen $1, 2, \dots$, sondern mit den Elementen $\alpha, \beta, \dots$ von $\mathcal{G}$ indizieren.

Man sieht sofort, daß der O. U_α , der durch

$$U_\alpha \left(\sum_{\beta \in \mathcal{G}} x_\beta \varphi_\beta \right) = \sum_{\beta \in \mathcal{G}} x_{\alpha\beta} \varphi_\beta \quad \left(\sum_{\beta \in \mathcal{G}} |x_\beta|^2 \text{ endlich!} \right)$$

definiert wird, ein unitärer ist. Dabei gilt $U_\alpha U_\beta = U_{\alpha\beta}$, also insbesondere $U_\alpha^* = U_\alpha^{-1} = U_{\alpha^{-1}}$. Sei nun $\mathcal{P}$ die Menge aller U_α . Wann gehört ein A von $\mathcal{B}$ zu $\mathcal{P}'$? A muß mit jedem U_α von $\mathcal{P}$ (die $U_\alpha^* = U_{\alpha^{-1}}$ liefern nichts Neues!) vert. sein, $A U_\alpha = U_\alpha A$. Dies bedeutet: $(A U_\alpha f, g) = (U_\alpha A f, g)$ für alle f, g — da aber beide Seiten linear und stetig in f und in g sind, genügt es, dies für ein vollst. norm. orth. System zu fordern. Z. B. für alle $f = \varphi_\beta$, $g = \varphi_\gamma$, wegen

$$(A U_\alpha \varphi_\beta, \varphi_\gamma) = (A \varphi_{\alpha^{-1}\beta}, \varphi_\gamma),$$

$$(U_\alpha A \varphi_\beta, \varphi_\gamma) = (A \varphi_\beta, U_\alpha^* \varphi_\gamma) = (A \varphi_\beta, U_{\alpha^{-1}} \varphi_\gamma) = (A \varphi_\beta, \varphi_{\alpha\gamma})$$

bedeutet das $(A \varphi_{\alpha^{-1}\beta}, \varphi_\gamma) = (A \varphi_\beta, \varphi_{\alpha\gamma})$ — d. h. daß $(A \varphi_\beta, \varphi_\gamma)$ nur von $\beta^{-1}\gamma$ abhängt.

Wenn also A, B zu $\mathcal{P}$ gehören, so ist (wir setzen $(A\varphi_\varepsilon, \varphi_\eta) = a(\varepsilon^{-1}\eta)$, $(B\varphi_\varepsilon, \varphi_\eta) = b(\varepsilon^{-1}\eta)$)

$$\begin{aligned} (AB\varphi_\beta, \varphi_\gamma) &= (B\varphi_\beta, A^*\varphi_\gamma) = \sum_{\delta \in \mathcal{G}} (B\varphi_\beta, \varphi_\delta) (\varphi_\delta, A^*\varphi_\gamma) \\ &= \sum_{\delta \in \mathcal{G}} (B\varphi_\beta, \varphi_\delta) (A\varphi_\delta, \varphi_\gamma) = \sum_{\delta \in \mathcal{G}} a(\delta^{-1}\gamma) b(\beta^{-1}\delta) = \sum_{\varepsilon, \eta \in \mathcal{G}, \varepsilon\eta = \beta^{-1}\gamma} a(\varepsilon) b(\eta), \end{aligned}$$

und genau so kommt

$$(BA\varphi_\beta, \varphi_\gamma) = \sum_{\varepsilon, \eta \in \mathcal{G}, \varepsilon\eta = \beta^{-1}\gamma} a(\varepsilon) b(\eta)$$

heraus. Daher gilt $(ABf, g) = (BAf, g)$ zumindest im vollst. norm. orth. System der φ_α , also, da beide Seiten linear und stetig in f und in g sind, für alle f, g . Somit ist $AB = BA$, A, B sind vert., und damit ist $\mathcal{P}'$ (welches ja mit A auch A^* enthält) als Abelsch erkannt.

2. Auf die soeben konstruierte Menge $\mathcal{P}$ wollen wir nun das am Anfang des § 1 Gesagte anwenden, d. h. die R, R^* dieser Klasse untersuchen.

Sei a_α eine Zahlenfolge (gleichfalls mit Elementen α von $\mathcal{G}$ indiziert), wobei nur die Endlichkeit von $\sum_{\alpha \in \mathcal{G}} |a_\alpha|^2$ vorausgesetzt werde. Wir definieren zwei $O. R, R^*$ für die φ_α , und nur diese, folgendermaßen:

$$R\varphi_\alpha = \sum_{\beta \in \mathcal{G}} a_{\alpha^{-1}\beta} \varphi_\beta, \quad R^*\varphi_\alpha = \sum_{\beta \in \mathcal{G}} \overline{a_{\alpha\beta^{-1}}} \varphi_\beta$$

($\sum_{\beta \in \mathcal{G}} |a_{\alpha^{-1}\beta}|^2 = \sum_{\beta \in \mathcal{G}} |a_{\alpha\beta^{-1}}|^2 = \sum_{\beta \in \mathcal{G}} |a_\beta|^2$ ist endlich). Man sieht sofort, daß R, R^* im Sinne von Definition 5 ein konj. Paar bilden; ferner, daß R mit allen U_α (also auch den $U_\alpha^* = U_{\alpha^{-1}}$) vert. ist — d. h. $(R)'\supset \mathcal{P}$ gilt. Also gehört R, R^* zur Klasse $\mathcal{P}$; es ist normal, wir können daher $\tilde{R}, \tilde{R}^*$ bilden. Dieses ist max. normal und gehört auch zur Klasse $\mathcal{P}$.

Nun bestimmen wir, auf Grund von Satz 15, den Definitionsbereich und den Wertverlauf von $\tilde{R}$ und $\tilde{R}^*$. Danach sind $\tilde{R}f, \tilde{R}^*f$ dann und nur dann sinnvoll, wenn zwei f^*, f^{**} existieren, so daß (soweit Rg, R^*g Sinn haben)

$$(f, Rg) = (f^{**}, g), \quad (f, R^*g) = (f^*, g)$$

gilt. Da $g = \varphi_\alpha$ sein muß, bedeutet dies, wenn wir noch $f = \sum_{\beta \in \mathcal{G}} x_\beta \varphi_\beta$ setzen ($\sum_{\beta \in \mathcal{G}} |x_\beta|^2$ endlich):

$$\sum_{\beta \in \mathcal{G}} x_\beta \overline{a_{\alpha^{-1}\beta}} = (f^{**}, \varphi_\alpha), \quad \sum_{\beta \in \mathcal{G}} x_\beta a_{\alpha\beta^{-1}} = (f^*, \varphi_\alpha).$$

Wir setzen

$$y_\alpha = \sum_{\beta \in \mathcal{G}} a_{\alpha\beta^{-1}} x_\beta, \quad z_\alpha = \sum_{\beta \in \mathcal{G}} \overline{a_{\alpha^{-1}\beta}} x_\beta,$$

dann muß also $(f^*, \varphi_\alpha) = y_\alpha$, $(f^{**}, \varphi_\alpha) = z_\alpha$ sein — nach E. § I, Satz 5, 7 ist dies dann und nur dann möglich, wenn $\sum_{\alpha \in \mathcal{G}} |y_\alpha|^2$, $\sum_{\alpha \in \mathcal{G}} |z_\alpha|^2$ endlich sind, und zwar ist dann $f^* = \sum_{\alpha \in \mathcal{G}} y_\alpha \varphi_\alpha$, $f^{**} = \sum_{\alpha \in \mathcal{G}} z_\alpha \varphi_\alpha$. Unser Resultat ist also:

$\tilde{R}f$, $\tilde{R}^*f$ haben für $f = \sum_{\alpha \in \mathcal{G}} x_\alpha \varphi_\alpha$ ($\sum_{\alpha \in \mathcal{G}} |x_\alpha|^2$ endlich) dann und nur dann Sinn, wenn für die soeben definierten y_α , z_α die Summen $\sum_{\alpha \in \mathcal{G}} |y_\alpha|^2$, $\sum_{\alpha \in \mathcal{G}} |z_\alpha|^2$ endlich sind, und zwar ist dann $Rf = \sum_{\alpha \in \mathcal{G}} y_\alpha \varphi_\alpha$, $R^*f = \sum_{\alpha \in \mathcal{G}} z_\alpha \varphi_\alpha$.

Andererseits besitzt aber $\tilde{R}$, $\tilde{R}^*$ nach Satz 17, 18 genau eine Z. d. E. $\Delta(z)$; wie in § 1 festgestellt wurde, gehören die $\Delta(z)$ zu $\mathcal{P}$. In unserem Falle bedeutet dies, wie wir ebendort sahen, daß $(\Delta(z) \varphi_\alpha, \varphi_\beta)$ nur von $\alpha^{-1} \beta$ abhängt: dieses ist das allgemeine Element der Matrix des P. O. $\Delta(z)$ im vollst. norm. orth. System $\varphi_\alpha, \varphi_\beta, \dots$ —

Wählen wir für $\mathcal{G}$ z. B. die unendliche zyklische Gruppe (z. B. alle ganzen Zahlen $0, \pm 1, \pm 2, \dots$ mit der Addition als Kompositionsregel), so gewinnen wir eine allgemeine Eigenwerttheorie der Laurentschen Matrizen⁷¹⁾ (auch der nicht-reellen und unbeschränkten). Andere Wahlen von $\mathcal{G}$ führen zu analogen, aber allgemeineren Matrizenklassen, die wir nun gleicherweise beherrschen.

Anhang III.

1. Hier soll der Vertauschbarkeitsbegriff für unbeschr. O. etwas näher analysiert werden. Wir hatten die Normalität von R , d. i. die Vert. von R, R^* , durch den Abelschen Charakter von $(R)''$ definiert — also wenn wir die H. O. $S_1 = \frac{R+R^*}{2}$, $S_2 = \frac{R-R^*}{2i}$ bilden, deren Vert. durch den Abelschen Charakter von $(S_1, S_2)''$.⁷²⁾ Es liegt aber nahe, zu diesem Zwecke die übliche Definition der Matrizenvert. heranzuziehen zu versuchen.

Seien z. B. R, S zwei H. O., $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System, in welchem beide eine Matrix besitzen: $\{a_{\mu\nu}\}$ bzw. $\{b_{\mu\nu}\}$.⁷³⁾ Da die Zeilen- und Spalten-Absolutwertquadratsummen $\sum_{\mu=1}^{\infty} |a_{\mu\nu}|^2 = \sum_{\mu=1}^{\infty} |a_{\nu\mu}|^2$, $\sum_{\mu=1}^{\infty} |b_{\mu\nu}|^2 = \sum_{\mu=1}^{\infty} |b_{\nu\mu}|^2$ alle endlich sind, konvergieren auch die Reihen

⁷¹⁾ Diese und ähnliche Klassen von Matrizen hat Toeplitz mit funktionentheoretischen Fragen in Beziehung gebracht. Vgl. auch a. a. O. Anm. ²⁴⁾.

⁷²⁾ Beim Beweise von Satz 14 sahen wir $(R)' = (R, R^*)'$, also $= (S_1, S_2)'$, $(R)'' = (S_1, S_2)''$.

⁷³⁾ Vgl. E., Anhang III, ferner die demnächst im Journal f. Math. erscheinende Arbeit des Verfassers „Zur Theorie der unbeschränkten Matrizen“.

$$\sum_{\varrho=1}^{\infty} a_{\mu\varrho} b_{\varrho\nu}, \quad \sum_{\varrho=1}^{\infty} b_{\mu\varrho} a_{\varrho\nu} \text{ absolut, und man könnte versuchen, die Vert. durch}$$

$$\sum_{\varrho=1}^{\infty} a_{\mu\varrho} b_{\varrho\nu} = \sum_{\varrho=1}^{\infty} b_{\mu\varrho} a_{\varrho\nu} \quad (\mu, \nu = 1, 2, \dots)$$

zu definieren.

Wegen $a_{\mu\nu} = (R\varphi_{\mu}, \varphi_{\nu})$, $b_{\mu\nu} = (S\varphi_{\mu}, \varphi_{\nu})$ (vgl. z. B. E., Anhang III) ist aber

$$\sum_{\varrho=1}^{\infty} a_{\mu\varrho} b_{\varrho\nu} = \sum_{\varrho=1}^{\infty} (R\varphi_{\mu}, \varphi_{\varrho}) (S\varphi_{\varrho}, \varphi_{\nu}) = \sum_{\varrho=1}^{\infty} (R\varphi_{\mu}, \varphi_{\varrho}) (\varphi_{\varrho}, S\varphi_{\nu}) = (R\varphi_{\mu}, S\varphi_{\nu}),$$

$$\sum_{\varrho=1}^{\infty} b_{\mu\varrho} a_{\varrho\nu} = \sum_{\varrho=1}^{\infty} (S\varphi_{\mu}, \varphi_{\varrho}) (R\varphi_{\varrho}, \varphi_{\nu}) = \sum_{\varrho=1}^{\infty} (S\varphi_{\mu}, \varphi_{\varrho}) (\varphi_{\varrho}, R\varphi_{\nu}) = (S\varphi_{\mu}, R\varphi_{\nu}),$$

also lautet die Vert.-Bedingung einfach so:

$$(R\varphi_{\mu}, S\varphi_{\nu}) = (S\varphi_{\mu}, R\varphi_{\nu}).$$

Wenn nun alle $R(S\varphi_{\mu})$, $S(R\varphi_{\mu})$ sinnvoll sind, so können wir hierfür schreiben:

$$(SR\varphi_{\mu}, \varphi_{\nu}) = (RS\varphi_{\mu}, \varphi_{\nu}), \quad ((RS - SR)\varphi_{\mu}, \varphi_{\nu}) = 0,$$

und da die φ_{ν} ein vollst. norm. orth. System bilden, $(RS - SR)\varphi_{\mu} = 0$. Wenn also irgendein abg. lin. O. T Forts. von $i(RS - SR)$ ist, so gilt $Tf = 0$ für alle $f = \varphi_{\mu}$, also auch für alle Linearaggregate endlich vieler — d. h. in einer überall dichten Menge. Da nun T abg. ist, ist danach Tf stets sinnvoll, und zwar $= 0$, also $T = 0$. Andererseits ist $i(RS - SR)$ offenbar H. (es ist für alle φ_{μ} sinnvoll, d. h. sein Definitionsbereich spannt die abg. lin. M. $\mathfrak{H}$ auf), somit können wir $T = \overline{i(RS - SR)}$ bilden, und dieses muß $= 0$ sein.

In diesem Falle können wir also mit einigem Recht von einer Vert. von R, S reden — insbesondere für R, S aus $\mathcal{B}$ (beschr.!) trifft dies zu, da aber $i(RS - SR)$ dann sogar überall sinnvoll, also selbst abg. lin. ist, ist $i(RS - SR) = 0$, $RS = SR$. Wie steht es aber bei beliebigen R, S , für die die $R(S\varphi_{\mu})$, $S(R\varphi_{\mu})$ nicht alle sinnvoll sein müssen? Inwieweit bedeutet dann $(R\varphi_{\mu}, S\varphi_{\nu}) = (S\varphi_{\mu}, R\varphi_{\nu})$ ($\mu, \nu = 1, 2, \dots$, im vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$ mögen R, S Matrizen haben; vgl. Anm. ⁷⁸) eine vernünftige Vert. von R, S ?

Im folgenden soll an einem Gegenbeispiel gezeigt werden, daß die genannte Relation in dieser Allgemeinheit kaum von Bedeutung sein kann.

Zunächst sei aber noch erwähnt, auf welchen Begriff der Normalität (für konj. Paare R, R^*) dieser Vertauschbarkeitsbegriff (für H. O.) führt.

Wir müssen dann $S_1 = \frac{R+R^*}{2}$, $S_2 = \frac{R-R^*}{2i}$ bilden, und

$$(S_1\varphi_{\mu}, S_2\varphi_{\nu}) = (S_2\varphi_{\mu}, S_1\varphi_{\nu})$$

verlangen, oder, wie man leicht umrechnet,

$$(R \varphi_\mu, R \varphi_\nu) = (R^* \varphi_\mu, R^* \varphi_\nu).$$

Dies bedeutet offenbar:

$$|Rf| = |R^*f|$$

für alle Linearaggregate endlich vieler φ_μ . Daß das für die Normalität notwendig ist, wissen wir aus Satz 21.

2. Sei $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System. Wir definieren zwei Matrizen $\{a_{\mu\nu}\}, \{b_{\mu\nu}\}$ so: wenn nicht μ und ν gleich $2n-1$ oder $2n$ (für dasselbe $n=1, 2, \dots$) ist, so sei $a_{\mu\nu} = b_{\mu\nu} = 0$; ferner

$$a_{2n-1, 2n-1} = a_{2n, 2n-1} = a_{2n-1, 2n} = a_{2n, 2n} = n;$$

$$b_{2n-1, 2n-1} = b_{2n, 2n} = n, \quad b_{2n-1, 2n} = -b_{2n, 2n-1} = in.$$

Die zu den (offenbar H. quadrierbaren) Matrizen $\{a_{\mu\nu}\}$ bzw. $\{b_{\mu\nu}\}$ (und $\{\varphi_1, \varphi_2, \dots\}$) gehörigen abg. lin. H. O. seien R bzw. S (vgl. Anm. ⁷³).

Durch

$$\psi_{2n-1} = \frac{1}{\sqrt{2}} \varphi_{2n-1} + \frac{1}{\sqrt{2}} \varphi_{2n}, \quad \psi_{2n} = \frac{1}{\sqrt{2}} \varphi_{2n-1} - \frac{1}{\sqrt{2}} \varphi_{2n};$$

$$\chi_{2n-1} = \frac{1}{\sqrt{2}} \varphi_{2n-1} + \frac{i}{\sqrt{2}} \varphi_{2n}, \quad \chi_{2n} = \frac{1}{\sqrt{2}} \varphi_{2n} - \frac{i}{\sqrt{2}} \varphi_{2n-1}$$

sind ebenfalls vollst. norm. Orth. Systeme $\psi_1, \psi_2, \dots$ und $\chi_1, \chi_2, \dots$ definiert, und zwar gilt

$$R\psi_{2n-1} = 2n \cdot \psi_{2n-1}, \quad R\psi_{2n} = 0; \quad S\chi_{2n-1} = 2n \cdot \chi_{2n-1}, \quad S\chi_{2n} = 0,$$

dabei sind die R, S abg. lin. In diesen vollst. norm. orth. Systemen haben also R bzw. S Diagonalmatrizen, also sind beide hypermax.⁷⁴). Insbesondere hat Rf für $f = \sum_{\nu=1}^{\infty} y_\nu \psi_\nu$ ($\sum_{\nu=1}^{\infty} |y_\nu|^2$ endlich) dann und nur dann Sinn, wenn $\sum_{n=1}^{\infty} 4n^2 \cdot |y_{2n-1}|^2$ endlich ist; für $f = \sum_{\nu=1}^{\infty} x_\nu \varphi_\nu$ ($\sum_{\nu=1}^{\infty} |x_\nu|^2$ endlich) ist aber $y_{2n-1} = \frac{1}{\sqrt{2}} x_{2n-1} + \frac{1}{\sqrt{2}} x_{2n}$, daher ist die Endlichkeit von $\sum_{n=1}^{\infty} n^2 |x_{2n-1} + x_{2n}|^2$ charakteristisch. Ebenso zeigt man, daß die Sinnvollheit von Sf für $f = \sum_{\nu=1}^{\infty} z_\nu \chi_\nu$ ($\sum_{\nu=1}^{\infty} |z_\nu|^2$ endlich) durch die Endlichkeit von $\sum_{n=1}^{\infty} 4n^2 |z_{2n-1}|^2$, also für $f = \sum_{\nu=1}^{\infty} x_\nu \varphi_\nu$, wo $z_{2n-1} = \frac{1}{\sqrt{2}} x_{2n-1} - \frac{i}{\sqrt{2}} x_{2n}$ ist, durch diejenige von $\sum_{n=1}^{\infty} n^2 |x_{2n-1} - ix_{2n}|^2$ gekennzeichnet ist. Die Endlichkeit beider Summen ist, wie man leicht einsieht, mit der von

⁷⁴) Allgemein folgt aus $T\xi_n = \alpha_n \xi_n$ ($n=1, 2, \dots$, T ein abg. lin. O., $\xi_1, \xi_2, \dots$ ein vollst. norm. orth. System) die Hypermax. von T . In der Tat umfassen die abg. lin. M. $\mathfrak{E}, \mathfrak{F}$ (vgl. E. Kap. V) alle $(T \pm i1)\xi_n = (\alpha_n \pm i)\xi_n$, also alle ξ_n . Daher sind beide gleich $\mathfrak{F}$.

$\sum_{n=1}^{\infty} n^2 (|x_{2n-1}|^2 + |x_{2n}|^2)$ gleichwertig. Wenn nun noch $R(Sf)$ und $S(Rf)$ sinnvoll sein sollen, so ist $Rf = \sum_{v=1}^{\infty} u_v \varphi_v$, $Sf = \sum_{v=1}^{\infty} v_v \varphi_v$ mit

$$u_{2n-1} = n \cdot x_{2n-1} + n \cdot x_{2n}, \quad u_{2n} = n \cdot x_{2n-1} + n \cdot x_{2n};$$

$$v_{2n-1} = n \cdot x_{2n-1} - i n \cdot x_{2n}, \quad v_{2n} = i n \cdot x_{2n-1} + n \cdot x_{2n}$$

zu betrachten: es muß $\sum_{n=1}^{\infty} n^2 |v_{2n-1} + v_{2n}|^2$ und $\sum_{n=1}^{\infty} n^2 |u_{2n-1} - i u_{2n}|^2$ endlich sein, d. h. $\sum_{n=1}^{\infty} n^4 |x_{2n-1} + x_{2n}|^2$ und $\sum_{n=1}^{\infty} n^4 |x_{2n-1} - i x_{2n}|^2$. Man überlegt sich leicht, daß die Endlichkeit dieser Summen derjenigen von $\sum_{n=1}^{\infty} n^4 (|x_{2n-1}|^2 + |x_{2n}|^2)$ gleichwertig ist.

Nun bestimmen wir für diese f $i(RS - SR)f = \sum_{v=1}^{\infty} w_v \varphi_v$, es ergibt sich:

$$z_{2n-1} = -2n^2 \cdot x_{2n-1}, \quad z_{2n} = 2n^2 \cdot x_{2n}.$$

Daher ist dieses $i(RS - SR)$ ein abg. lin. H. O., und hat eine Diagonalmatrix im System $\varphi_1, \varphi_2, \dots$ — also ist es hypermax. (vgl. Anm. ⁷⁴). Daß es $\neq 0$ ist, ist klar: wie man sieht, folgt sogar aus $f \neq 0$ $i(RS - SR)f \neq 0$.

Von einer Vert. von R, S kann also keine Rede sein. Trotzdem werden wir ein vollst. norm. orth. System $\omega_1, \omega_2, \dots$ auffinden, in dem R, S beide Matrizen besitzen, und die Vert.-Relationen des § 1 erfüllt sind.

3. Daß R und S in einem vollst. norm. orth. System $\omega_1, \omega_2, \dots$ Matrizen $\{a_{\mu v}\}$ bzw. $\{b_{\mu v}\}$ besitzen, besagt folgendes (vgl. Anm. ⁷³): es muß jedenfalls $a_{\mu v} = (R\omega_\mu, \omega_v)$, $b_{\mu v} = (S\omega_\mu, \omega_v)$ sein, die elementaren O. dieser Matrizen, R' bzw. S' , sind die auf die Menge der $\omega_1, \omega_2, \dots$ eingeschränkten O. R bzw. S , die abg. lin. O. sind $\tilde{R}'$ bzw. $\tilde{S}'$ — und es wird verlangt $\tilde{R}' = R$, $\tilde{S}' = S$. D. h.: zu jedem f mit sinnvollem Rf bzw. Sf soll es eine Folge $f_1, f_2, \dots$ aus der lin. M. (nicht der abg. lin. M.!) der $\omega_1, \omega_2, \dots$ geben, mit $f_n \rightarrow f$ und $Rf_n \rightarrow Rf$ bzw. $Sf_n \rightarrow Sf$ (wenn Rf, Sf beide Sinn haben, können es zwei verschiedene Folgen sein!). Wenn $\omega_1, \omega_2, \dots$ durch das E. Schmidtsche Orthogonalisationsverfahren ⁷⁵) aus einer überall dichten Folge $g_1, g_2, \dots$ entsteht, so ist dies gewiß der Fall, wenn die soeben genannten $f_1, f_2, \dots$ mit $f_n \rightarrow f$ und $Rf_n \rightarrow Rf$ bzw. $Sf_n \rightarrow Sf$ als Teilfolge von $g_1, g_2, \dots$ gewählt werden können. Und wenn stets $(Rg_p, Sg_q) = (Sg_p, Rg_q)$ gilt, so gilt dies auch für die $\omega_1, \omega_2, \dots$ (die Linearaggregate der g sind): $(R\omega_\mu, S\omega_\nu) = (S\omega_\mu, R\omega_\nu)$, d. h. R, S erfüllen die Vert.-Relation des § 1. Es sei noch betont, daß wir die Sinnvollheit der Rg_p, Sg_p voraussetzten, aber nicht die der $R(Sg_p), S(Rg_p)$.

⁷⁵) Vgl. E. Schmidt, am in Anm. ⁵⁹) a. O. Siehe auch E. § 1, Satz 8.

Sei $\bar{g}_1, \bar{g}_2, \dots$ die Menge aller Linearaggregate endlich vieler $\varphi_1, \varphi_2, \dots$ mit rationalen Koeffizienten, als Folge geschrieben. Da R, S im System $\varphi_1, \varphi_2, \dots$ Matrizen haben, sieht man sofort, daß die früher erwähnten Folgen $f_1, f_2, \dots$ aus $\bar{g}_1, \bar{g}_2, \dots$ ausgewählt werden können (die Rationalitätsbedingung für die Koeffizienten stört nicht). Seien nun $h_1, h_2, \dots$ so gewählt, daß alle Rh_n, Sh_n Sinn haben, und (wie immer in $\S$, stark)

$$h_n \rightarrow 0, \quad Rh_{2n-1} \rightarrow 0, \quad Sh_{2n} \rightarrow 0 \quad (\text{für } n \rightarrow \infty)$$

gilt. Dann sieht man, daß die obigen Folgen $f_1, f_2, \dots$ stets auch als Teilfolgen von $\bar{g}_1 + h_1, \bar{g}_2 + h_2, \dots, \bar{g}_1 + h_2, \bar{g}_2 + h_4, \dots$ gewählt werden können. Wir schreiben für $\bar{g}_1, \bar{g}_2, \bar{g}_3, \bar{g}_4, \dots$ kürzer $\tilde{g}_1, \tilde{g}_2, \dots$, und wählen $g_1, g_2, \dots$ als $\tilde{g}_1 + h_1, \tilde{g}_2 + h_2, \dots$. Daher ist alles erreicht, wenn noch $(Rg_p, Sg_q) = (Sg_p, Rg_q)$ gilt — dies gilt es durch geeignete Wahl der $h_1, h_2, \dots$ zu erzwingen (die Wahl der $h_1, h_2, \dots$ ist ja noch bis zu einem gewissen Grade frei, während die $\tilde{g}_1, \tilde{g}_2, \dots$ bereits festgelegt sind).

Wir indizieren die Folge $\varphi_1, \varphi_2, \dots$ folgendermaßen um: Wir teilen sie in Paare $\varphi_{2n-1}, \varphi_{2n}$ ein und ersetzen die einfache Indizierung $n = 1, 2, \dots$ durch eine dreifache mit p, q, μ (alle $= 1, 2, \dots$), und zwar sollen $\varphi_{2n-1}, \varphi_{2n}$ bzw. $\varphi'_{pq\mu}, \varphi''_{pq\mu}$ heißen. Das zu p, q, μ gehörige n heiße $n(pq\mu)$, wir wollen es so einrichten, daß stets $n(pq\mu) \geq \mu^2$ ist. Nun wird h_p definiert

$$h_p = \sum_{q\mu=1}^{\infty} \frac{1}{pq \cdot (n(pq\mu))^{\frac{3}{2}}} \cdot \varphi'_{pq\mu} + \sum_{r=1}^{p-1} x_{rp} \cdot v_{rp},$$

dabei ist

$$v_{rp} = \begin{cases} \varphi'_{r,p,e(rp)} - \varphi''_{r,p,e(rp)}, & \text{für ungerades } p \\ \varphi'_{r,p,e(rp)} - i \varphi''_{r,p,e(rp)}, & \text{für gerades } p \end{cases}.$$

Über die x_{rp} und $e(rp)$ ($r < p$) verfügen wir noch später. Betrachten

wir zuerst die Summen $\xi_p = \sum_{q,\mu=1}^{\infty} \frac{1}{pq \cdot (n(pq\mu))^{\frac{3}{2}}} \varphi'_{pq\mu}$. Sie sind konvergent, ihre Absolutwertquadratsummen sind

$$\sum_{q,\mu=1}^{\infty} \frac{1}{p^2 q^2 (n(pq\mu))^3} \leq \sum_{q,\mu=1}^{\infty} \frac{1}{p^2 q^2 \mu^6} = \frac{1}{p^2} \left(\sum_{q=1}^{\infty} \frac{1}{q^2} \right) \left(\sum_{\mu=1}^{\infty} \frac{1}{\mu^6} \right),$$

daher streben sie mit $p \rightarrow \infty$ gegen 0. Da auch die Summen

$$\begin{aligned} \sum_{q,\mu=1}^{\infty} \frac{1}{p^2 q^2 (n(pq\mu))^3} (n(pq\mu))^2 &= \sum_{q,\mu=1}^{\infty} \frac{1}{p^2 q^2 n(pq\mu)} \leq \sum_{q,\mu=1}^{\infty} \frac{1}{p^2 q^2 \mu^2} \\ &= \frac{1}{p^2} \left(\sum_{q=1}^{\infty} \frac{1}{q^2} \right) \left(\sum_{\mu=1}^{\infty} \frac{1}{\mu^2} \right) \end{aligned}$$

endlich sind (sie streben sogar mit $p \rightarrow \infty$ gegen 0), sind auch die $R\xi_p, S\xi_p$ sinnvoll. $\sum_{r=1}^{p-1} x_{rp} \cdot v_{rp}$ ist eine endliche Summe, daher ist darauf R und S anwendbar; sie strebt mit $p \rightarrow \infty$ gegen 0, wenn $\sum_{r=1}^{p-1} |x_{rp}|^2 \rightarrow 0$, was z. B. für $|x_{rp}| \leq \frac{1}{rp}$ gewiß der Fall ist. Ferner ist für ungerades bzw. gerades p $Rv_{rp} = 0$ bzw. $Sv_{rp} = 0$, d. h. R bzw. S auf $\sum_{r=1}^{p-1} x_{rp} v_{rp}$ angewandt ergibt 0. Zusammenfassend sehen wir: Rh_p, Sh_p sind stets sinnvoll, $h_p \rightarrow \infty$, $Rh_p \rightarrow 0$ für ungerade p , $Sh_p \rightarrow 0$ für gerade p .

Es gilt nun noch $(Rg_p, Sg_q) - (Sg_p, Rg_q) = 0$ durch geeignete Wahl der $x_{rp}, \varrho(rp)$ zu sichern (man sieht, daß es genügt, nur $p < q$ zu betrachten, was geschehen soll) — natürlich unter Wahrung von $|x_{rp}| \leq \frac{1}{rp}$. Wir setzen in die obige Formel

$$g_p = \tilde{g}_p + \sum_{t, \mu=1}^{\infty} \frac{1}{p t (n(p t \mu))^{\frac{3}{2}}} \cdot \varphi'_{p t \mu} + \sum_{r=1}^{p-1} x_{rp} \cdot v_{rp},$$

$$g_q = \tilde{g}_q + \sum_{\mu, r=1}^{\infty} \frac{1}{q u (n(q u r))^{\frac{3}{2}}} \cdot \varphi'_{q u r} + \sum_{s=1}^{q-1} x_{sq} \cdot v_{sq}$$

ein, und rechnen den Ausdruck aus. Es wird, unter Berücksichtigung der bestehenden Orthogonalitäten (es ist $p < q!$)

$$\begin{aligned} & (\{SR - RS\} \tilde{g}_p, \tilde{g}_q) + (\{SR - RS\} \tilde{g}_p, \xi_q) + (\xi_p, \{RS - SR\} \tilde{g}_q) \\ & + \sum_{s=1}^{q-1} \bar{x}_{sq} (\{SR - RS\} \tilde{g}_p, v_{sq}) + \sum_{r=1}^{p-1} x_{rp} (v_{rp}, \{RS - SR\} \tilde{g}_q) \\ & + \frac{\bar{x}_{pq}}{pq(n(p, q, \varrho(pq)))^{\frac{3}{2}}} (\{SR - RS\} \varphi'_{p, q, \varrho(pq)}, v_{pq}) = 0. \end{aligned}$$

Wir setzen zur Abkürzung $i(RS - SR) = T$. Ferner führen wir die folgende Bezeichnung ein: $\tilde{g}_i$ ist ein Linearaggregat endlich vieler φ_n , also endlich vieler $\varphi'_{pq\mu}$ und $\varphi''_{pq\mu}$, das größte dabei vorkommende μ sei $\mu(i)$. Wir schreiben nun vor, daß stets $\varrho(rt) > \text{Max}(\mu(1), \dots, \mu(t-1))$ sein soll — dann fällt in unserem Ausdruck das vierte Glied identisch fort. Es wird:

$$\begin{aligned} & (T\tilde{g}_p, \tilde{g}_q) + (T\tilde{g}_p, \xi_q) + (\xi_p, T\tilde{g}_q) + \sum_{r=1}^{p-1} x_{rp} (v_{rp}, T\tilde{g}_q) \\ & + \frac{\bar{x}_{pq}}{pq(n(p, q, \varrho(pq)))^{\frac{3}{2}}} (T\varphi'_{p, q, \varrho(pq)}, v_{pq}) = 0. \end{aligned}$$

Wenn wir noch $(T\varphi'_{p, q, \varrho(pq)}, v_{pq}) = -2(n(p, q, \varrho(pq)))^2$ beachten, so wird

$$\begin{aligned} \bar{x}_{pq} = & - \frac{1}{2(n(p, q, \varrho(pq)))^{\frac{3}{2}}} \left[(T\tilde{g}_p, \tilde{g}_q) + (T\tilde{g}_p, \xi_q) + (\xi_p, T\tilde{g}_q) \right. \\ & \left. + \sum_{r=1}^{p-1} x_{rp} (v_{rp}, T\tilde{g}_q) \right]. \end{aligned}$$

Hier ist x_{pq} mit Hilfe der x_{rp} und v_{rp} , also der $\varrho(rp)$, mit $r = 1, \dots, p-1$ sowie von $\varrho(pq)$ ausgedrückt. Dabei ist $\varrho(pq)$ willkürlich. Wenn wir also die x_{rs} , $\varrho(rs)$ nach wachsendem $r+s$ anordnen, so liegt hier eine Rekursionsgleichung für die x_{rs} vor — und die $\varrho(rs)$ sind dabei beliebig. Es gilt nur noch $|x_{pq}| \leq \frac{1}{pq}$ und $\varrho(pq) \geq \text{Max}(\mu(1), \dots, \mu(q-1))$ einzuhalten. Wir wählen die $\varrho(rs)$ auch in der genannten Reihenfolge, dann haben wir (C ist der Absolutwert des Inhalts von [...]) in der Formel für $\bar{x}_{pq}$, hängt also nur von x_{rs} , $\varrho(rs)$ mit $r+s < p+q$ ab, von $\varrho(pq)$ aber nicht!)

$$|x_{pq}| = \frac{C}{2(n(p, q, \varrho(pq)))^{\frac{1}{2}}} \leq \frac{C}{2\varrho(pq)}.$$

Somit genügt es, $\varrho(pq) \geq \frac{1}{2}Cpq$ und $\geq \text{Max}(\mu(1), \dots, \mu(q-1))$ zu wählen, und wir sind am Ziele.

Wir haben jetzt die Folgen $h_1, h_2, \dots$ sowie $g_1, g_2, \dots$ und das System $\omega_1, \omega_2, \dots$, auf das es uns ankam: damit ist die gewünschte Konstruktion durchgeführt.

Zur Theorie der unbeschränkten Matrizen.

Einleitung.

I. In einer früheren Arbeit ¹⁾ wurde die Theorie der Hermiteschen Operatoren (kurz: H. O.) des Hilbertschen Raumes behandelt, und unabhängig von jeder Beschränktheits- oder Regularitäts-Annahme die Diskussion ihres Eigenwertproblems durchgeführt. In bezug auf die hierbei erzielten Resultate sei auf die genannte Arbeit verwiesen, auf die wir uns hier auch sonst stützen werden. Der Gegenstand der vorliegenden Arbeit ist, die Theorie von den Operatoren auf die Matrizen zu übertragen.

Das Verhältnis Operator-Matrix — das in endlichvielfdimensionalen (Euklidischen) Räumen, und auch im Hilbertschen Raume bei beschränkten Operatoren, ein einfaches und zwar ein ein-eindeutiges ist — weist nämlich bei unbeschränkten Operatoren des Hilbertschen Raumes neue Züge auf. Es ist dort wesentlich komplizierter, und die Theorie der Matrizen ist keineswegs, wie in den obengenannten Fällen, mit der der Operatoren gleichbedeutend und gleichlautend. Der Umstand an dem dies liegt, läßt sich wohl am besten dadurch kennzeichnen, daß man im Endlichvielfdimensionalen, und auch im Hilbertschen Raume bei beschränkten Operatoren, einem H. O. R und einem willkürlichen vollständigen normierten orthogonalen System (kurz: vollst. norm. orth. System — bezüglich dieser und aller folgenden Begriffsbildungen vgl. etwa A. E. Kap. I) $\varphi_1, \varphi_2, \dots$ eine Matrix $A = \{a_{\mu\nu}\}$ zuordnen konnte:

$$a_{\mu\nu} = (\varphi_\mu, R\varphi_\nu) \quad (\mu, \nu = 1, 2, \dots).$$

Bei einem unbeschränkten Operator R aber, der nach einem Satze von *Toeplitz* nicht überall sinnvoll sein kann (vgl. z. B. A. E. Satz 48 und Anm. 61), ist diese Bildung überhaupt nur dann möglich, wenn alle $R\varphi_1, R\varphi_2, \dots$ sinnvoll sind. Und auch wenn dies der Fall ist, treten in der Transformationstheorie dieser Matrizen eine Reihe von Konvergenzschwierigkeiten und Beschränkungen auf, die der „beschränkten“ Theorie völlig fremd sind.

So kommt es, daß z. B. die allgemein übliche Zuordnung Operator—Matrix nicht einmal unitär-invariant ist (vgl. Satz 3, 4). Infolgedessen kann, obwohl bei den Operatoren alles recht vernünftig ist (vgl. die Resultate von A. E.), bei den Matrizen eine eigentümliche Pathologie bestehen. Wie weit diese geht, ist aus dieser Arbeit zu ersehen, vgl. hauptsächlich Satz 6, 8, 9, 10, 14. Es gelingt aber trotzdem, eine ziemlich gute Übersicht über die hier herrschenden Verhältnisse zu gewinnen.

II. Zu erwähnen ist noch, daß während bei Operatoren in A. E. überhaupt keine Beschränktheits-ähnlichen Annahmen gemacht werden mußten, hier immerhin eine (wenn auch recht schwache) Bedingung zu stellen ist: die Endlichkeit aller Zeilen-Ab-

¹⁾ Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren, Math. Annalen Bd. 102/1, Seite 49—131. Sie soll als „A. E.“ zitiert werden.

solutwertquadratsummen $\sum_1^\infty |a_{\mu\nu}|^2$ ($\mu = 1, 2, \dots$; — natürlich dürfen beliebig große Zahlen darunter vorkommen). Dies liegt daran, daß einem ganz willkürlichen H. O. R (und einem vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$ mit lauter sinnvollen $R\varphi_1, R\varphi_2, \dots$) nicht eine beliebige Hermitesche Matrix $A = \{a_{\mu\nu}\}$ ($a_{\mu\nu} = \overline{a_{\nu\mu}}$) entspricht, sondern eine mit lauter endlichen $\sum_1^\infty |a_{\mu\nu}|^2$ — das wurde in A. E. Anhang III gezeigt. Diese

Matrizen, wir werden sie Hermitesch quadrierbar (kurz: H. quadr.) nennen, bilden also den Gegenstand unserer Untersuchungen. Ihre Theorie wurde bereits im genannten Anhang III in Angriff genommen; an diesen werden wir im folgenden anknüpfen. Die für uns wichtigsten Resultate und Definitionen werden wir zwar wiederholen, im allgemeinen müssen wir aber die Terminologie und die Resultate von A. E. benützen.

III. Schließlich sei noch darauf hingewiesen, daß unsere Untersuchungen in enger Beziehung zur neuen Quantentheorie stehen, zur sog. „Matrizentheorie“ und „Transformationstheorie“ — deren mathematischer Apparat in erster Linie von den unitären Transformationen unbeschränkter (aber in allen Fällen quadrierbarer) Hermitescher Matrizen gebildet wird ²⁾ — was gerade der Gegenstand dieser Arbeit ist. Es zeigt sich indessen, daß von den einfachen Verhältnissen, die bei endlichvieldeimensionalen, sowie bei unendlichvieldeimensionalen aber beschränkten Matrizen bestehen — und deren Fortbestehen auch im Unbeschränkten für diese Theorien höchst wesentlich zu sein scheint — nur sehr wenig übrig bleibt. Bei Operatoren ist wohl das meiste in Ordnung (vgl. A. E.): das Eigenwertproblem hat keine oder genau eine Lösung, und wenn es keine hat, so treten an Stelle der Hilbertschen „Spektralform“ andere einfache Normalformen, usw. Bei Matrizen aber herrscht, wie in I. angedeutet wurde, eine so unerwartete Pathologie, daß ein Aufbau auf dieser Grundlage recht schwer zu sein scheint. Das eigentümlichste dabei ist, daß sich die Hermiteschen Matrizen verhältnismäßig vernünftig benehmen, die eigentliche Quelle der Pathologie sind die unitären Transformationen (von denen dies, wegen ihrer Beschränktheit, eigentlich nicht zu erwarten war). Wir werden dies im folgenden noch genauer auseinandersetzen.

I. Die unitär-konvergente Äquivalenz.

1. In A. E. Anhang III, nannten wir eine unendliche Matrix $A = \{a_{\mu\nu}\}$ ($\mu, \nu = 1, 2, \dots$) H. quadr., wenn $a_{\mu\nu} = \overline{a_{\nu\mu}}$ gilt, und alle $\sum_1^\infty |a_{\mu\nu}|^2$ endlich sind. Wenn $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System ist, so wurde durch $R\varphi_\mu = \sum_1^\infty a_{\mu\nu} \varphi_\nu$ (Rf sonst sinnlos!) der elementare Operator von $A, \{\varphi_1, \varphi_2, \dots\}$ definiert — die H. O. $\hat{R}$ und $\tilde{R}$ (vgl. die diesbezüglichen Ausführungen in A. E., insbesondere Satz 9 und Anhang III 1.) waren der lineare bzw. der abgeschlossene lineare H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$.

Der Definitionsbereich von $\hat{R}$ sind alle $\sum_1^N x_\mu \varphi_\mu$ ($N = 1, 2, \dots$), der von $\tilde{R}$ ist schwerer zu charakterisieren. Dagegen waren die Erweiterungselemente $f = \sum_1^\infty x_\mu \varphi_\mu$ und die ihnen zugeordneten f^* (vgl. A. E. Anhang III 1.) leicht angebbar: man bilde

²⁾ Zusammenfassende Darstellungen dieser Theorie gaben *Jordan*, Zeitschr. f. Phys. Bd. 37, Seite 383, und *Dirac*, Proc. Royal Soc. Bd. 118 A, Seite 621. Vgl. auch das Buch von *Weyl*, Gruppentheorie und Quantenmechanik, Leipzig 1928.

$$y_\nu = \sum_1^\infty a_{\mu\nu} x_\mu$$

(alle Reihen konvergieren absolut), f ist Erweiterungselement wenn $\sum_1^\infty |y_\nu|^2$ endlich ist, und dann $f^* = \sum_1^\infty y_\nu \varphi_\nu$. (Dies ist insbesondere der Wert von $\widehat{R}f$ oder $\widetilde{R}f$, wenn diese Sinn haben.)

Ferner wurde gezeigt, daß zu jedem abgeschlossenen linearen H. O. (kurz: abg. lin. H. O.) R eine H. quadr. Matrix $A = \{a_{\mu\nu}\}$ und ein vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$ gefunden werden kann, so daß R der abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ ist. D. h.: die abg. lin. H. O. aller $A, \{\varphi_1, \varphi_2, \dots\}$ machen gerade die Gesamtheit aller abg. lin. H. O. aus.

Die Zuordnung ist aber ein-mehrdeutig in dem Sinne, daß ein abg. lin. H. O. R zu mehreren $A, \{\varphi_1, \varphi_2, \dots\}$ gehören kann. Dies wollen wir jetzt näher untersuchen.

2. Wenn R und $\varphi_1, \varphi_2, \dots$ gegeben sind, so ist dadurch A bestimmt; denn R soll der abg. lin. H. O. von A sein, also Fortsetzung seines elementaren H. O. S , also

$$S\varphi_\mu = \sum_1^\infty a_{\mu\nu} \varphi_\nu, \quad a_{\mu\nu} = (S\varphi_\mu, \varphi_\nu) = (R\varphi_\mu, \varphi_\nu).$$

Wenn wir aber $\varphi_1, \varphi_2, \dots$ irgendwie (mit sinnvollen $R\varphi_1, R\varphi_2, \dots$!) wählen, und $A = \{a_{\mu\nu}\}$ mit $a_{\mu\nu} = (R\varphi_\mu, \varphi_\nu)$ bilden (es ist H. quadr., vgl. A. E. Anhang III) und die zu $A, \{\varphi_1, \varphi_2, \dots\}$ gehörigen Operatoren $S, \widehat{S}, \widetilde{S}$ (elementar, lin., abg. lin.), so ist jedenfalls

$$S\varphi_\mu = \sum_1^\infty a_{\mu\nu} \varphi_\nu = \sum_1^\infty (R\varphi_\mu, \varphi_\nu) \cdot \varphi_\nu = R\varphi_\mu,$$

also R Forts. von S , und weil es abg. lin. ist, auch von $\widehat{S}, \widetilde{S}$ — aber $R = \widetilde{S}$ ist zunächst fraglich.

Die Verhältnisse liegen vielmehr so:

Satz 1. Sei R ein abg. lin. H. O., $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System. Wenn es eine (H. quadr.) Matrix $A = \{a_{\mu\nu}\}$ gibt, so daß R der abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ ist, so müssen jedenfalls alle $R\varphi_1, R\varphi_2, \dots$ Sinn haben, und es muß $a_{\mu\nu} = (\varphi_\mu, R\varphi_\nu)$ sein — d. h. es gibt zu gegebenen $\varphi_1, \varphi_2, \dots$ kein oder genau ein A .

Wenn R beschränkt ist, gibt es wirklich zu allen $\varphi_1, \varphi_2, \dots$ ein solches A ; ist es aber unbeschränkt, so gibt es immer sowohl solche $\varphi_1, \varphi_2, \dots$ für die A existiert, als auch solche, für die das nicht der Fall ist (wobei sogar alle $R\varphi_1, R\varphi_2, \dots$ sinnvoll sind!).

Beweis: Die erste Gruppe von Behauptungen besteht aus schon vorher bewiesenen Tatsachen, wir müssen also die zweite betrachten.

Sei R beschränkt. Es ist überall sinnvoll (A. E. Satz 48, α, γ), wir können zu $\varphi_1, \varphi_2, \dots$ das A und dessen Operatoren $S, \widehat{S}, \widetilde{S}$ bilden. Wir sahen, daß R Forts. des abg. lin. H. O. S ist, da R beschränkt ist, muß $\widetilde{S} = R$ sein (A. E. Satz 48, α, δ)³⁾. — Sei R unbeschränkt. Daß $\varphi_1, \varphi_2, \dots$ so gewählt werden kann, daß A existiert, wissen wir schon; es gilt die andere Möglichkeit zu exemplifizieren. R ist sicher eig. Forts.

³⁾ Bei beiden Anwendungen des ziemlich tiefen Satzes 48 kommt es nur auf die triviale Hinreichendheit seiner Beschränktheitskriterien an.

eines abg. lin. H. O. T (A. E. Satz 48, α, δ), dieser ist sicher der abg. lin. H. O. geeigneter $A, \{\varphi_1, \varphi_2, \dots\}$. Da alle $T\varphi_1, T\varphi_2, \dots$ sinnvoll sind, sind es auch alle $R\varphi_1, R\varphi_2, \dots$, und es ist

$$A = \{a_{\mu\nu}\}, \quad a_{\mu\nu} = (\varphi_\mu, T\varphi_\nu) = (\varphi_\mu, R\varphi_\nu);$$

trotzdem ist der dazugehörige abg. lin. H. O. $T \neq R$.

3. Schon in A. E. Anhang III haben wir erörtert, wie zwei $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$, die denselben abg. lin. H. O. haben, zusammenhängen. Um den Gegenstand wieder aufnehmen zu können, definieren wir zuerst:

Definition 1. $A = \{a_{\mu\nu}\}$ sei eine H. quadr. Matrix, $U = \{u_{\mu\nu}\}$ eine unitäre Matrix d. h. es soll

$$\sum_1^\infty u_{\mu\varrho} \overline{u_{\nu\varrho}} = \begin{cases} 1 & \text{für } \mu = \nu, \\ 0 & \text{für } \mu \neq \nu, \end{cases} \quad \sum_1^\infty u_{\varrho\mu} \overline{u_{\varrho\nu}} = \begin{cases} 1 & \text{für } \mu = \nu, \\ 0 & \text{für } \mu \neq \nu \end{cases}$$

gelten. Wir nennen U auf A konvergent anwendbar, wenn die folgenden Bedingungen erfüllt sind:

α . Die Reihen $\sum_1^\infty a_{\varrho\sigma} u_{\sigma\mu}, \sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}}$ konvergieren sowieso absolut⁴⁾, außerdem sollen auch die Summen $\sum_1^\infty \left| \sum_1^\infty a_{\varrho\sigma} u_{\sigma\nu} \right|^2, \sum_1^\infty \left| \sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} \right|^2$ endlich sein (wenn es die eine ist, ist es auch die andere⁵⁾).

β . Infolgedessen konvergieren die Reihen $\sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right), \sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right)$ absolut⁶⁾, sie sollen dieselbe Summe haben (d. h. $\sum_1^\infty, \sum_1^\sigma$ sollen vertauschbar sein).

Den gemeinsamen Wert der beiden Summen nennen wir $b_{\mu\nu}$:

$$b_{\mu\nu} = \sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right) = \sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right).$$

$B = \{b_{\mu\nu}\}$ ist dann die Transformierte von A . —

Die Matrix B ist wieder H. quadr.: $b_{\mu\nu} = \overline{b_{\nu\mu}}$ folgt sofort aus der Definition, und die Endlichkeit von $\sum_1^\infty |b_{\mu\nu}|^2 = \sum_1^\infty \left| \sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} \right) u_{\sigma\nu} \right|^2$ folgt aus der von $\sum_1^\infty \left| \sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} \right|^2$ und der Unitarität von $\{u_{\mu\nu}\}$ ⁷⁾.

⁴⁾ Weil $\sum_1^\infty |a_{\varrho\sigma}|^2, \sum_1^\infty |u_{\sigma\mu}|^2 = 1, \sum_1^\infty |a_{\varrho\sigma}|^2 = \sum_1^\infty |a_{\sigma\varrho}|^2, \sum_1^\infty |u_{\varrho\mu}|^2 = 1$ alle endlich sind.

⁵⁾ Denn es ist $\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} = \overline{\sum_1^\infty a_{\sigma\varrho} u_{\varrho\mu}}$.

⁶⁾ Wir können sie in der Form $\sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} u_{\sigma\nu} \right) \overline{u_{\varrho\mu}}, \sum_1^\infty \left(\sum_1^\infty a_{\varrho\sigma} \overline{u_{\varrho\mu}} \right) u_{\sigma\nu}$ schreiben.

⁷⁾ Wenn $\sum_1^\infty |x_\sigma|^2$ endlich ist, so sei $f = \sum_1^\infty x_\sigma \varphi_\sigma$, dann ist

$$\begin{aligned} \sum_1^\infty x_\sigma u_{\sigma\nu} &= \sum_1^\infty (f, \varphi_\sigma) (\varphi_\sigma, \psi_\nu) = (f, \psi_\nu), \\ \sum_1^\infty \left| \sum_1^\infty x_\sigma u_{\sigma\nu} \right|^2 &= \sum_1^\infty |(f, \psi_\nu)|^2 = \|f\|^2 = \sum_1^\infty |x_\sigma|^2. \end{aligned}$$

Definition 2. $A = \{a_{\mu\nu}\}$ sei eine H. quadr. Matrix, $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System. Wenn $U = \{u_{\mu\nu}\}$ eine auf A konvergent anwendbare unitäre Matrix ist, so verstehen wir unter der Anwendung von U auf $A, \{\varphi_1, \varphi_2, \dots\}$ das Folgende: Man bilde die H. quadr. Matrix B nach Def. 1 und das vollst. norm. orth. System der $\psi_1, \psi_2, \dots, \psi_\mu = \sum_{\nu=1}^{\infty} \overline{u_{\nu\mu}} \varphi_\nu$ *) — dann ist $B, \{\psi_1, \psi_2, \dots\}$ das transformierte System.

Zwei willkürliche $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ heißen konvergent-unitär-äquivalent (kurz: k. u. äquivalent, natürlich sollen A, B H. quadr. sein, und $\varphi_1, \varphi_2, \dots$ und $\psi_1, \psi_2, \dots$ vollst. norm. orth.), wenn eine unitäre Matrix, die dieses bewirkt, existiert 9).

Wir werden im folgenden $A, \{\varphi_1, \varphi_2, \dots\}$ alle Eigenschaften seines abg. lin. H. O. zuschreiben — d. h. es max., hypermax. (A. E. Def. 7, 9), definit, beschränkt, halb-beschränkt (A. E. Def. 12), Forts. oder eig. Forts. eines anderen solchen Systems (A. E., vor Satz 9), usw. nennen —, wenn sein abg. lin. H. O. die betreffende Eigenschaft hat. Max., Hypermax., Definität, Beschränktheit und beide Halbbeschränktheiten kommen sogar A an und für sich zu; denn wenn ein $A, \{\varphi_1, \varphi_2, \dots\}$ sie besitzt, so besitzen sie alle $A, \{\psi_1, \psi_2, \dots\}$, da der Übergang vom einen zum anderen durch einen Isomorphismus von $\mathfrak{S}$ bewirkt werden kann, bei dem diese Eigenschaften invariant sind.

Satz 2 10). A ist dann und nur dann beschränkt, wenn alle unitären U auf A konvergent anwendbar sind.

Beweis: Sei $\varphi_1, \varphi_2, \dots$ irgendein vollst. norm. orth. System. Wenn U auf A konvergent anwendbar ist, so führt es $A, \{\varphi_1, \varphi_2, \dots\}$ in $B, \{\psi_1, \psi_2, \dots\}$ über; dann hat der abg. lin. H. O. R von $A, \{\varphi_1, \varphi_2, \dots\}$ auch zu $\{\psi_1, \psi_2, \dots\}$ ein B , so daß er zu $B, \{\psi_1, \psi_2, \dots\}$ gehört. Ist umgekehrt dies der Fall, so sind $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ k. u. äquivalent (nach dem gleich zu beweisenden Satz 3); da aber bereits $\varphi_1, \varphi_2, \dots$ und $\psi_1, \psi_2, \dots$ die Transformationsmatrix bestimmen, muß diese gleich U sein — also ist U auf A konvergent anwendbar.

*) Die Reihen. $\sum_{\nu=1}^{\infty} \overline{u_{\nu\mu}} \varphi_\nu$ konvergieren, weil $\sum_{\nu=1}^{\infty} |u_{\nu\mu}|^2 = 1$ endlich ist. Die ψ_μ sind norm. orth. wegen

$$(\psi_\mu, \psi_\nu) = \sum_{\varrho=1}^{\infty} (\psi_\mu, \varphi_\varrho) (\overline{\varphi_\nu, \varphi_\varrho}) = \sum_{\varrho=1}^{\infty} \overline{u_{\varrho\mu}} u_{\varrho\nu} = \begin{cases} 1, & \text{für } \mu = \nu, \\ 0, & \text{für } \mu \neq \nu. \end{cases}$$

Da weiter

$$\begin{aligned} |\varphi_\mu - \sum_{\nu=1}^N u_{\nu\mu} \psi_\nu|^2 &= |\varphi_\mu|^2 + \sum_{\nu=1}^N |u_{\nu\mu}|^2 |\psi_\nu|^2 - 2\Re \sum_{\nu=1}^N \overline{u_{\nu\mu}} (\varphi_\mu, \psi_\nu) \\ &= 1 + \sum_{\nu=1}^N |u_{\nu\mu}|^2 - 2\Re \sum_{\nu=1}^N |u_{\nu\mu}|^2 = 1 - \sum_{\nu=1}^N |u_{\nu\mu}|^2 \end{aligned}$$

ist, also mit $N \rightarrow \infty$ gegen 0 strebt, ist $\sum_{\nu=1}^{\infty} u_{\nu\mu} \psi_\nu$ konvergent und $= \varphi_\mu$. Die abg. lin. M. der $\psi_1, \psi_2, \dots$ enthält also alle $\varphi_1, \varphi_2, \dots$, sie ist somit $= \mathfrak{S}$; die $\psi_1, \psi_2, \dots$ sind also vollständig (A. E. Satz 7, α). (Der Beweis ist im wesentlichen nach E. Schmidt, Rend. d. Circ. Math. d. Palermo 25, 1908.)

9) Wegen $\psi_\mu = \sum_{\nu=1}^{\infty} \overline{u_{\nu\mu}} \varphi_\nu$ ist dieses U festgelegt: $\overline{u_{\nu\mu}} = (\psi_\mu, \varphi_\nu)$, $u_{\mu\nu} = (\varphi_\mu, \psi_\nu)$. Das so bestimmte U ist allenfalls unitär und transformiert $\varphi_1, \varphi_2, \dots$ in $\psi_1, \psi_2, \dots$ (vgl. A. E. Anhang III), die Frage ist also: erstens ob U auf A konvergent anwendbar ist, und zweitens ob dabei B entsteht.

10) Dieser Satz steht in engem Zusammenhange mit gewissen Untersuchungen von Toeplitz, Math. Ann. Bd. 69, Seite 211.

Unsere Bedingung bedeutet also: zu jedem $\{\psi_1, \psi_2, \dots\}$ existiert ein B , so daß R der abg. lin. H. O. von B , $\{\psi_1, \psi_2, \dots\}$ ist, und nach Satz 1 ist dies für die Beschränktheit von R (also von A) charakteristisch.

4. Anstatt die am Anfang von § 2 erwähnten Überlegungen aus A. E. Anhang III zu wiederholen, beweisen wir gleich eine Verschärfung.

Satz 3. $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ sind dann und nur dann k. u. äquivalent, wenn sie eine gemeinsame Forts. (einen H. O. ¹¹⁾) haben.

Beweis: Die elementaren, lin., abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ seien $R, \tilde{R}, \tilde{R}$ bzw. $S, \tilde{S}, \tilde{S}$. Die einzige in Frage kommende Transformationsmatrix ist

$$U = \{u_{\mu\nu}\}, \quad u_{\mu\nu} = (\varphi_\mu, \psi_\nu),$$

und die k. u. Äquivalenz bedeutet: U ist auf $A = \{a_{\mu\nu}\}$ konvergent anwendbar und es entsteht dabei $B = \{b_{\mu\nu}\}$ (vgl. Anm. 9).

Sei erstens dies der Fall. Dann ist

$$\begin{aligned} (R\varphi_\mu, \psi_\nu) &= \sum_1^\infty (R\varphi_\mu, \varphi_\epsilon) (\varphi_\epsilon, \psi_\nu) = \sum_1^\infty a_{\mu\epsilon} u_{\nu\epsilon}, \\ (\varphi_\mu, S\psi_\nu) &= \sum_1^\infty (\varphi_\mu, \psi_\epsilon) (\psi_\epsilon, S\psi_\nu) = \sum_1^\infty b_{\epsilon\nu} u_{\mu\epsilon}. \end{aligned}$$

Aus Def. 1 folgt

$$\sum_1^\infty b_{\epsilon\nu} u_{\mu\epsilon} = \sum_1^\infty \left[\sum_1^\infty \left\{ \sum_1^\infty a_{\sigma\tau} u_{\tau\nu} \right\} \overline{u_{\sigma\epsilon}} \right] u_{\mu\epsilon},$$

dies ist aber gleich $\sum_1^\infty a_{\mu\tau} u_{\tau\nu}$ ¹²⁾; also ist $(R\varphi_\mu, \psi_\nu) = (\varphi_\mu, S\psi_\nu)$. Hieraus folgt sofort: wenn $\tilde{R}f, \tilde{S}g$ Sinn haben, so ist $(\tilde{R}f, g) = (f, \tilde{S}g)$; und wenn $\tilde{R}f, \tilde{S}g$ Sinn haben, so ist $(\tilde{R}f, g) = (f, \tilde{S}g)$.

Wenn $\tilde{R}f$ Sinn hat, so ist demnach f ein Erweiterungselement von $\tilde{S}$, u. zw. ist ihm $\tilde{R}f$ zugeordnet; wenn also auch $\tilde{S}f$ Sinn hat, so ist $\tilde{R}f = \tilde{S}f$. Wir können also einen Operator T so definieren: Tf hat Sinn, wenn $\tilde{R}f$ oder $\tilde{S}f$ Sinn hat, und ist dann diesem gleich. Da $\tilde{R}, \tilde{S}$ H. O. sind und die obige Relation gilt, ist auch T ein H. O. Außerdem ist es offenbar Forts. von $\tilde{R}, \tilde{S}$.

Existiere zweitens ein H. O. T , der Forts. von $\tilde{R}$ und $\tilde{S}$ ist. Wir schließen, wie in A. E. Anhang III:

$$\sum_1^\infty a_{\epsilon\sigma} u_{\sigma\nu} = \sum_1^\infty (\tilde{R}\varphi_\epsilon, \varphi_\sigma) (\varphi_\sigma, \psi_\nu) = (\tilde{R}\varphi_\epsilon, \psi_\nu) = (T\varphi_\epsilon, \psi_\nu) = (\varphi_\epsilon, T\psi_\nu) = (\varphi_\epsilon, \tilde{S}\psi_\nu),$$

so daß $\sum_1^\infty \left| \sum_1^\infty a_{\epsilon\sigma} u_{\sigma\nu} \right|^2 = \sum_1^\infty |(\varphi_\epsilon, \tilde{S}\psi_\nu)|^2 = |\tilde{S}\psi_\nu|^2$ endlich ist, und weiter:

¹¹⁾ Da dieser H. O. durch jede Forts. von sich ersetzt werden kann, dürfen wir ihn gleich max. annehmen (A. E. Satz 35).

¹²⁾ Sei $\sum_1^\infty |x_\mu|^2$ endlich, dann setzen wir $f = \sum_1^\infty \overline{x_\mu} \varphi_\mu$. Es ist

$$\sum_1^\infty \left[\sum_1^\infty x_\sigma \overline{u_{\sigma\epsilon}} \right] u_{\mu\epsilon} = \sum_1^\infty \left[\sum_1^\infty (\varphi_\sigma, f) (\psi_\epsilon, \varphi_\sigma) \right] (\varphi_\mu, \psi_\epsilon) = \sum_1^\infty (\psi_\epsilon, f) (\varphi_\mu, \psi_\epsilon) = (\varphi_\mu, f) = x_\mu.$$

$$\sum_1^{\infty} \left(\sum_1^{\infty} a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right) = \sum_1^{\infty} \left(\sum_1^{\infty} a_{\varrho\sigma} u_{\sigma\nu} \right) \overline{u_{\varrho\mu}} = \sum_1^{\infty} (\varphi_{\varrho}, \tilde{S} \psi_{\nu}) (\psi_{\mu}, \varphi_{\varrho}) = (\psi_{\mu}, \tilde{S}_{\nu} \psi_{\nu}) = b_{\mu\nu}.$$

Wenn wir hier μ, ν (und ϱ, σ) vertauschen, die konjugiert-komplexen nehmen, und $a_{\varrho\sigma} = \overline{a_{\sigma\varrho}}$, $b_{\mu\nu} = \overline{b_{\nu\mu}}$ beachten, so wird schließlich:

$$b_{\mu\nu} = \sum_1^{\infty} \left(\sum_1^{\infty} a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \right) = \sum_1^{\infty} \left(\sum_1^{\infty} a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\varrho\nu} \right).$$

Also ist U auf A konvergent anwendbar, und führt es in B über.

Das Bestehen dieses Satzes zeigt, daß die Verallgemeinerung des Hilbertschen Begriffes der unitären Äquivalenz beschränkter Matrizen, d. h. die k. u. Äquivalenz, so wie wir sie in Def. 1, 2 getroffen haben, schwerlich auf wesentlich anderer Grundlage gelingen kann. Denn wir haben in Def. 1 an Konvergenzen eher zu viel als zu wenig verlangt — unsere diesbezüglichen Forderungen sind wohl kaum zu erhöhen¹³⁾ — und erreichten trotzdem, 1. daß bei beschränkten Operatoren alles beim alten blieb, 2. daß Matrizen (und vollst. norm. orth. Systeme) mit demselben Operator k. u. äquivalent sind.

Dennoch erweist Satz 3 auch solche Matrizen, die nicht denselben Operator haben, manchmal als k. u. äquivalent. Aus diesem Umstande folgt der größte Teil der Matrizen-Pathologie.

II. Folgerungen.

1. Ob $A, \{\varphi_1, \varphi_2, \dots\}$ dem $B, \{\psi_1, \psi_2, \dots\}$ k. u. äquivalent ist, hängt nach Satz 3 nur von ihren abg. lin. H. O. ab, die R bzw. S seien — wir können also die Eigenschaft der k. u. Äquivalenz R und S selbst zuschreiben. Wir führen übrigens dafür das Zeichen $\sim$ ein. (R, S müssen abg. lin. H. O. sein!)

In erster Linie müssen wir feststellen, wieweit $\sim$ die typischen Eigenschaften der Äquivalenzen (identische, reflexive und transitive, u. ä.) zukommen.

Satz 4. Die max. H. O. können dadurch charakterisiert werden, daß sie nur ihren Teilen k. u. äquivalent sind; die beschränkten dadurch, daß sie nur sich selbst k. u. äquivalent sind. Zwei k. u. äquivalente max. H. O. sind identisch.

Beweis: Sei R max. Wenn $R \sim S$ ist, so gibt es eine gemeinsame Forts. T von R, S ; da R max. ist, ist $R = T$; also ist R Forts. von S . Sei umgekehrt R nicht max. Es gibt eine max. Forts. S von R ; da R, S die gemeinsame Forts. S haben, ist $R \sim S$. Wäre R Forts. von S , so müßte $R = S$ sein, d. h. R wäre max., entgegen der Annahme.

Sei R beschränkt. Aus $R \sim S$ folgt, da R max. ist (vgl. hier und für die folgenden Schlüsse A. E. Satz 48), S Teil von R . Aber R hat keine anderen abg. lin. H. O. Teile als sich selbst, also $R = S$. Sei umgekehrt R unbeschränkt. Dann ist es eig. Forts. eines abg. lin. H. O. S ; da R, S die gemeinsame Forts. R haben, ist $R \sim S$, dabei $R \neq S$.

Die dritte Behauptung folgt aus der ersten.

¹³⁾ Man könnte daran denken, analog wie in der Hilbertschen Theorie (bei beschränkten Formen)

$$\sum_1^N \sigma a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \rightarrow b_{\mu\nu} \quad (N \rightarrow \infty) \quad \text{oder sogar} \quad \sum_{1,1}^{M,N} \sigma a_{\varrho\sigma} \overline{u_{\varrho\mu}} u_{\sigma\nu} \rightarrow b_{\mu\nu} \quad (M, N \rightarrow \infty)$$

zu verlangen. Indessen läßt sich zeigen: zu jedem unbeschränkten $A, \{\varphi_1, \varphi_2, \dots\}$ gibt es ein $B, \{\psi_1, \psi_2, \dots\}$ mit demselben abg. lin. H. O., das ihm im Sinne keiner dieser zwei Bedingungen unitär-äquivalent ist.

uns an die Berg- und Tal-Skizzen in §§ 3, 4 von Kap. II (besonders vor Satz 7) erinnern, so sehen wir, daß dies bedeutet, daß R, USU^{-1} in drei Schritten k. u. äquivalent sind. Dabei ist R der Operator von $A, \{\varphi_1, \varphi_2, \dots\}$ und USU^{-1} der von $B, \{U\varphi_1, U\varphi_2, \dots\}$, also sind nach Def. 3 die Matrizen A, B in drei Schritten k. u. äquivalent.

Es bleibt noch übrig, zwei A, B zu finden, die nicht in zwei Schritten k. u. äquivalent sind. Sei A max. und nach unten aber nicht nach oben halbbeschränkt, und B max. und nach oben aber nicht nach unten halbbeschränkt²⁰⁾. Wären sie in zwei Schritten k. u. äquivalent, so gäbe es zwei Systeme $\varphi_1, \varphi_2, \dots$ und $\psi_1, \psi_2, \dots$, so daß für die abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$, R und S , k. u. Äquivalenz in zwei Schritten besteht. Nun sind R, S auch max., also müßten sie benachbart sein (vgl. § 4 in Kap. II); sei also T der abg. lin. H. O. von dem beide Forts. sind.

T muß wie R nach unten und wie S nach oben beschränkt sein, d. h. überhaupt beschränkt. Also ist es max. (A. E. Satz 48, α, β) und daher $R = S = T$, d. h. R, S beschränkt entgegen der Annahme.

7. Der soeben bewiesene Satz 10 läßt klar erkennen, wie es mit der k. u. Äquivalenz in mehreren Schritten bei H. quadr. Matrizen steht. Wenn wir sie in Klassen einteilen, so daß jede Klasse alle miteinander in endlich vielen Schritten äquivalenten H. quadr. Matrizen umfaßt, so bilden alle unbeschränkten eine einzige Klasse. Bei beschränkten ist es anders: Jede Klasse mit einem beschränkten Element A besteht nur aus beschränkten²¹⁾ Matrizen, und diese haben (bei geeigneter Wahl ihrer $\{\varphi_1, \varphi_2, \dots\}$) alle denselben abg. lin. H. O.; wenn wir also $\{\varphi_1, \varphi_2, \dots\}$ für alle Matrizen der Schar gemeinsam wählen, so bilden ihre abg. lin. H. O. eine Schar URU^{-1} , wo R ein fester beschränkter abg. lin. H. O. ist, und U alle unitären Operatoren durchläuft. Und zwar sind offenbar alle k. u. äquivalent (in einem Schritt, nicht nur in endlich vielen!).

Nebenbei sehen wir: Wenn zwei willkürliche Matrizen in endlich viel Schritten k. u. äquivalent sind, so sind sie es schon in drei Schritten — in zweien aber nicht immer.

IV. Die Äquivalenz aller unbeschränkten Operatoren.

1. Satz 10 sagte aus: Wenn A, B zwei willkürliche nicht beschränkte H. quadr. Matrizen sind, so sind sie in drei Schritten k. u. äquivalent, d. h. zu jedem $\{\varphi_1, \varphi_2, \dots\}$ gibt es ein $\{\psi_1, \psi_2, \dots\}$, so daß $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ in drei Schritten k. u. äquivalent sind. Wir werden diese Paradoxie weiter verschärfen, indem wir zeigen, daß auch noch $\{\psi_1, \psi_2, \dots\}$ vorgeschrieben werden kann; das bedeutet aber: zwei beliebige nicht beschränkte abg. lin. H. O. R, S sind in endlich vielen Schritten k. u. äquivalent²²⁾. Allerdings müssen wir die Zahl der Schritte von drei auf neun erhöhen, während wir die in Satz 10 gegebene (und hier erst recht gültige) untere Grenze für ihre Zahl, drei, nicht erhöhen können. Die in allen Fällen ausreichende Mindestzahl werden wir also nicht bestimmen können; sie ist ≥ 3 und ≤ 9 .

Wir beweisen zunächst eine Teilbehauptung:

Satz 11. Seien R, S die abg. lin. H. O. von $\{a_\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$ und $\{b_\mu \delta_{\mu\nu}\}, \{\psi_1, \psi_2, \dots\}$. Es gebe zwei elementfremde Teilfolgen $M_1, M_2, \dots$ und

²⁰⁾ Etwa beide Diagonalmatrizen $\{a_\mu \delta_{\mu\nu}\}$, und zwar bei A $a_\mu \rightarrow \infty$ und bei B $a_\mu \rightarrow -\infty$ (für $\mu \rightarrow \infty$).

²¹⁾ Seien A, B k. u. äquivalent, also $A, \{\varphi_1, \varphi_2, \dots\}, B, \{\psi_1, \psi_2, \dots\}$ ebenso. Wenn A , also der abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$, beschränkt ist, so muß nach Satz 4 $B, \{\psi_1, \psi_2, \dots\}$ denselben abg. lin. H. O. haben, also auch beschränkt sein. Dasselbe gilt also auch bei k. u. Äquivalenz in mehreren Schritten.

²²⁾ In Satz 10 war das Produkt der Transformierenden unitären Matrizen willkürlich, jetzt wird es auch noch beliebig vorgeschrieben.

Dieser Satz zeigt (wie in A. E. nach Satz 47 bemerkt wurde) daß die Abschnittsspektren einer unbeschränkten Matrix (bzw. die Menge ihrer Häufungspunkte) gar nicht die fundamentale Bedeutung haben, wie die einer beschränkten. Im letzteren Falle beherrschen sie ja die diesbezügliche Hilbertsche Theorie — bei uns sind sie, wie Satz 6 zeigt, nicht einmal unitär-invariant. Wir können z. B. eine weder nach oben noch nach unten halbbeschränkte H. quadr. und max. Matrix A nehmen; diese ist dann einem B mit lauter Abschnitts-Eigenwerten ≥ 1 und einem C mit lauter Abschnitts-Eigenwerten ≤ -1 k. u. äquivalent (Satz 6). — und da A max. ist, ist sogar B dem C k. u. äquivalent (Satz 5)!

Man beachte übrigens, daß unsere Sätze nicht etwa nur für einige besonders böswillig gewählte Matrizen gelten, sondern für alle unbeschränkten ohne Ausnahme!

3. Da die k. u. Äquivalenz sich als nicht transitiv erwiesen hat, ist es naheliegend, die folgenden Begriffe einzuführen:

Definition 3. $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ sollen in n Schritten k. u. äquivalent heißen, in Zeichen: $\sim_n$, wenn es $n+1$ solche Systeme $A^\nu, \{\varphi_1^\nu, \varphi_2^\nu, \dots\}$, ($\nu = 0, \dots, n$), gibt¹⁵), daß $A = A^0, \{\varphi_1, \varphi_2, \dots\} = \{\varphi_1^0, \varphi_2^0, \dots\}$, $B = A^n, \{\psi_1, \psi_2, \dots\} = \{\varphi_1^n, \varphi_2^n, \dots\}$ ist, und $A^{\nu-1}, \{\varphi_1^{\nu-1}, \varphi_2^{\nu-1}, \dots\}$ mit $A^\nu, \{\varphi_1^\nu, \varphi_2^\nu, \dots\}$ k. u. äquivalent ist ($\nu = 1, \dots, n$). Wenn dies für irgend ein $n = 1, 2, \dots$ geht, so heißen sie in endlich vielen Schritten äquivalent. —

Die Begriffe k. u. äquivalent, in n bzw. in endlich vielen Schritten k. u. äquivalent können auch auf A, B allein (ohne die $\{\varphi_1, \varphi_2, \dots\}, \{\psi_1, \psi_2, \dots\}$) Anwendung finden: dann sind immer nur diese zu transformieren. (Man kann dann zu jedem $\{\varphi_1, \varphi_2, \dots\}$ ein $\{\psi_1, \psi_2, \dots\}$ finden, so daß $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ in dieser Relation stehen.)

Ebenso wie wir es am Anfang von § 1 dieses Kapitels taten, können wir auch die k. u. Äquivalenz in mehreren Schritten als Eigenschaft der abg. lin. H. O. der $A, \{\varphi_1, \varphi_2, \dots\}$ und $B, \{\psi_1, \psi_2, \dots\}$ ansehen; also bei abg. lin. H. O. von solchen Äquivalenzen reden.

Außerdem wollen wir die Tatsache, daß S Forts. oder eig. Forts. von R ist durch eine Ordnungsbeziehung ausdrücken: $S \leq$ bzw. $< R$. Dieses $<$ hat in der Tat die Eigenschaften einer unvollständigen Ordnung¹⁶). Die Relationen $R \sim S, R \sim_n S$ lassen sich mit Hilfe von Satz 3 dann folgendermaßen umschreiben: $R \sim S$ bedeutet, daß es ein T mit $R \leq T \leq S$ gibt (T darf nach Anm. 11 max. angenommen werden, selbstverständlich soll es allenfalls abg. lin. sein); $R \sim_n S$ bedeutet daß es $2n-1$ abg. lin. H. O. $T_1, S_1, T_2, S_2, \dots, T_{n-1}, S_{n-1}, T_n$ mit

$$R \leq T_1 \leq S_1 \leq T_2 \leq S_2 \leq \dots \leq T_{n-1} \leq S_{n-1} \leq T_n \leq S$$

gibt. Wir erwähnen noch: in der $R \sim_n S$ ausdrückenden Kette $R \sim S_1 \sim S_2 \sim \dots \sim S_{n-1} \sim S$ dürfen nach Satz 5 alle max. Zwischenglieder fortgelassen werden, wenn also ihrer k drin vorkommen, so gilt $R \sim_k S$ (die Max. der $T_1, \dots, T_n$ dagegen ist nichtssagend, wir können sie sogar immer erzwängen!).

4. Neben der k. u. Äquivalenz wollen wir für abg. lin. H. O. noch eine andere Art der Verknüpfung definieren.

Definition 4. R, S heißen benachbart, wenn sie beide Forts. eines und desselben

¹⁵) A, B, A^ν , usw. sollen immer H. quadr. Matrizen sein, $\{\varphi_1, \varphi_2, \dots\}, \{\psi_1, \psi_2, \dots\}, \{\varphi_1^\nu, \varphi_2^\nu, \dots\}$, usw. vollst. norm. orth. Systeme.

¹⁶) D. h. es ist nie $R < R$, aus $R < S, S < T$ folgt $R < T$, aber es muß nicht immer einer der drei Fälle $R <, =, > S$ bestehen; $R \leq S$ bedeutet $R = S$ oder $R < S$.

H. O. sind. Sie heißen in n Schritten benachbart, wenn es $n + 1$ abg. lin. H. O. $T_0, T_1, \dots, T_n$ gibt, so daß $R = T_0, S = T_n$ ist, und $T_{\nu-1}$ mit T_ν benachbart ist ($\nu = 1, \dots, n$)¹⁷⁾. Wenn dies für irgendein $n = 1, 2, \dots$ geht, so heißen sie in endlich vielen Schritten benachbart.

Benachbarkeit in einem Schritt bedeutet also $R \geq T \leq S$ (wir können T durch $\tilde{T}$ ersetzen, es also gleich als abg. lin. voraussetzen), und die in n Schritten

$$R \geq S_1 \leq T_1 \geq S_2 \leq T_2 \geq \dots \geq S_{n-1} \leq T_{n-1} \geq S_n \leq S$$

(wieder sind die $S_1, \dots, S_n$ abg. lin., die $T_1, \dots, T_{n-1}$ sogar max. annehmbar).

Wir sehen also: Benachbarkeit bzw. k. u. Äquivalenz von R und S in n Schritten bedeuten, daß ein „Berg- und Tal-Weg“ von R nach S führt, der aus n Bergen und $n - 1$ Tälern bzw. $n - 1$ Bergen und n Tälern besteht — im ersten Falle mit Bergen, im zweiten mit Tälern beginnt und aufhört. Hieraus folgt aber: wenn R, S in n Schritten benachbart sind, so sind sie in $n + 1$ Schritten k. u. äquivalent (man läßt die Sohle des ersten und letzten Tales mit R bzw. S zusammenfallen), wenn sie in n Schritten k. u. äquivalent sind, so sind sie in $n + 1$ Schritten benachbart (man läßt die Spitze des ersten und letzten Berges mit R bzw. S zusammenfallen). Im letzteren Falle können wir, falls einer oder beide von R, S max. ist, sogar auf Benachbarkeit in n bzw. $n - 1$ Schritten schließen; denn dann muß ja R , oder S , oder beide ohnehin auf der Spitze des ersten bzw. letzten Berges liegen, so daß man gleich mit dem danebenliegenden Tale anfangen bzw. aufhören kann.

Wir beweisen noch einen Satz über Benachbarkeit.

Satz 7. Wenn R, S benachbart sind, so ist für alle f , für die beide Sinn haben, $Rf = Sf$. Der für diese f definierte Operator T , der dort mit ihnen übereinstimmt, ihr „Durchschnitt“, ist ein abg. lin. H. O. von dem beide Forts. sind.

Beweis: Sei $R \geq S' \leq S$. Wenn Rf, Sf Sinn haben, so ist für alle g für die $S'g$ sinnvoll ist,

$$(Rf, g) = (f, Rg) = (f, S'g) = (f, Sg) = (Sf, g), \quad (Rf, g) = (Sf, g).$$

Das gilt also auch in ihrer abg. lin. M., d. h. für alle g . Also ist $Rf = Sf$.

T ist offenbar abg. lin.; da sein Definitionsbereich den von S' umfaßt, spannt er als abg. lin. M. $\mathfrak{S}$ auf. $(Tf, g) = (f, Tg)$ gilt offenbar, also ist T ein H. O. (A. E. Def. 6). Daß R, S Forts. von T sind, ist klar.

III. Die Äquivalenz aller unbeschränkten Matrizen.

1. Aus dem Titel dieses Kapitels geht sein Ziel hervor. Um es zu erreichen, müssen wir zunächst einige vorbereitende Überlegungen anstellen.

Eine Diagonalmatrix ist eine $A = \{a_{\mu\nu}\}$, für die aus $\mu \neq \nu$ $a_{\mu\nu} = 0$ folgt, d. h. eine von folgender Form: $A = \{a_\mu \delta_{\mu\nu}\}$ ($\delta_{\mu\nu} = \begin{matrix} 1 & \text{für } \mu = \nu \\ 0 & \text{für } \mu \neq \nu \end{matrix}$) — sie ist H. quadr. wenn alle a_μ reell sind. Den abg. lin. H. O. $\tilde{R}$ von $A, \{\varphi_1, \varphi_2, \dots\}$ nennen wir einen Diagonaloperator.

Er ist hypermax.: nach A. E. Anhang III müssen wir zeigen, daß die $(\tilde{R} + i \cdot 1)\varphi_\mu$ sowohl als auch die $(\tilde{R} - i \cdot 1)\varphi_\mu$ ganz $\mathfrak{S}$ als abg. lin. M. aufspannen. In der Tat enthält ihre abg. lin. M. alle $(\tilde{R} \pm i \cdot 1)\varphi_\mu = (a_\mu \pm i)\varphi_\mu$, also alle φ_μ , daher ist sie in beiden Fällen = $\mathfrak{S}$. Daher hat $Rf, f = \sum_{\mu=1}^{\infty} x_\mu \varphi_\mu$, für alle seine Erweiterungselemente Sinn

¹⁷⁾ Wie bei Anm. 11 können die $T_1, \dots, T_{n-1}$ max. angenommen werden.

(A. E. Def. 9), und diese sind (vgl. A. E. Anhang III) durch die Endlichkeit von $\sum_1^\infty a_\mu^2 |x_\mu|^2$ charakterisiert. Rf ist dann (vgl. a. a. O.) gleich $\sum_1^\infty a_\mu x_\mu \cdot \varphi_\mu$. — Aus diesen Ausführungen folgt sofort, daß jede abg. lin. M., die von irgendwelchen der $\varphi_1, \varphi_2, \dots$ aufgespannt wird, $\tilde{R}$ reduziert.

Wenn für irgendeinen abg. lin. H. O. S die Relationen $S\varphi_\mu = a_\mu \varphi_\mu$ ($\mu = 1, 2, \dots$) gelten, so ist S Forts. des elementaren H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$, von R . Also ist $S = \tilde{S}$ Forts. von $\tilde{R}$, und da dieses max. ist, $S = \tilde{R}$.

Diese Diagonaloperatoren werden jetzt eine wichtige Rolle spielen.

2. Sei R ein nach oben nicht halbbeschränkter abg. lin. H. O. Nach A. E. Satz 46 sowie Anm. 60) gibt es eine überall dichte Folge $f_1, f_2, \dots$ mit den folgenden Eigenschaften: Sei $\mathfrak{M}_p$ die von $f_p, f_{p+1}, \dots$ aufgespannte lin. M. (nicht abg.!, sie ist überall dicht, weil es $f_p, f_{p+1}, \dots$ sind), dann ist Rf für alle f von $\mathfrak{M}_p$ sinnvoll, und es gilt für diese durchweg

$$(f, Rf) \geq p \cdot |f|^2 \quad (p = 1, 2, \dots).$$

Das auf $\mathfrak{M}_p$ eingeschränkte R ist also ein lin. H. O. R_p . Aus der obigen Ungleichung folgt auch

$$(f, \tilde{R}_p f) \geq p \cdot |f|^2,$$

also ist für $f \neq 0$ $\tilde{R}_p f \neq 0$ (wenn es Sinn hat), und für $f \neq g$ $\tilde{R}_p f \neq \tilde{R}_p g$. Somit hat $\tilde{R}_p$ eine Inverse $\tilde{R}_p^{-1}$, die im Wertevorrat $\mathfrak{N}_p$ von $\tilde{R}_p$ definiert ist. Sie ist, ebenso wie $\tilde{R}_p$, abg. lin. Die für $\tilde{R}_p$ gültige Ungleichung besagt:

$$(\tilde{R}_p^{-1} g, g) \geq p \cdot |\tilde{R}_p^{-1} g|^2.$$

Wegen $|(\tilde{R}_p^{-1} g, g)| \leq |\tilde{R}_p^{-1} g| \cdot |g|$ folgt aber hieraus:

$$|\tilde{R}_p^{-1} g| \leq \frac{1}{p} \cdot |g|.$$

Also ist $\tilde{R}_p^{-1}$ stetig; da es auch abg. ist, muß sein Definitionsbereich $\mathfrak{N}_p$ abg. sein — eine lin. M. ist er ohnehin. Für $p < q$ ist $\tilde{R}_p$ Forts. von $\tilde{R}_q$, $\tilde{R}_p^{-1}$ Forts. von $\tilde{R}_q^{-1}$, also umfaßt $\mathfrak{N}_p$ dann $\mathfrak{N}_q$. Somit bilden $\mathfrak{N}_1, \mathfrak{N}_2, \dots$ eine Folge von abg. lin. M., deren jede alle späteren umfaßt. Dabei ist der Wertevorrat von R_p , d. i. die von den $Rf_p, Rf_{p+1}, \dots$ aufgespannte lin. M., sicher in $\mathfrak{N}_p$ überall dicht, als abg. lin. M. muß daher $\mathfrak{N}_p$ die von diesen aufgespannte abg. lin. M. sein. Hieraus schließt man mühelos, daß $\mathfrak{N}_p - \mathfrak{N}_q$ ($p < q$) höchstens $(q - p)$ -dimensional ist.

Wären alle $\mathfrak{N}_p = \mathfrak{Q}$, so wären alle $\tilde{R}_p^{-1}$ miteinander identisch, es würde also $|\tilde{R}_1^{-1} f| \leq \frac{1}{p} |f|$ für alle $p = 1, 2, \dots$ gelten; daher wäre identisch $\tilde{R}_1^{-1} f = 0$, was mit dem Charakter von $\tilde{R}_1^{-1}$ als Inverse unverträglich ist. Sei also etwa $\mathfrak{N}_p \neq \mathfrak{Q}$, $\mathfrak{Q} - \mathfrak{N}_p = \mathfrak{L} \neq (0)$.

3. Jetzt stellen wir eine Überlegung im Sinne von A. E. Satz 42 an. Wenn f zu $\mathfrak{L}$ gehört, so ist $(f, \tilde{R}_p g) = 0$, d. h. f ist Erweiterungselement von $\tilde{R}_p$, und zwar ist ihm 0 zugeordnet. Wenn also $\tilde{R}_p f$ Sinn hat, so ist $\tilde{R}_p f = 0$, $f = 0$. D. h. $\mathfrak{L}$ und der Definitionsbereich von $\tilde{R}_p$ sind lin. unabh. lin. M. (A. E. Def. 2). Wir betrachten diejenigen f , für die die Zerlegung $f = g + h$, $\tilde{R}_p g$ sinnvoll, h von $\mathfrak{L}$, möglich ist (sie

ist dann eindeutig!) und definieren für diese einen Operator S durch $Sf = \tilde{R}_{\bar{p}}g$. S ist Forts. von $\tilde{R}_{\bar{p}}$ und ein H. O.; denn wenn Sf, Sf' Sinn haben ($f = g + h, f' = g' + h'$, $\tilde{R}_{\bar{p}}g, \tilde{R}_{\bar{p}}g'$ sinnvoll, h, h' von $\mathfrak{L}$), so ist:

$$(Sf, f') = (\tilde{R}_{\bar{p}}g, g' + h') = (\tilde{R}_{\bar{p}}g, g') = (g, \tilde{R}_{\bar{p}}g') = (g + h, \tilde{R}_{\bar{p}}g') = (f, Sf').$$

$\mathfrak{L}$ reduziert S (vgl. A. E. Def 13): in $\mathfrak{L}$ ist S immer sinnvoll, und zwar $= 0$, weiter gehört Sf (da es $\tilde{R}_{\bar{p}}g$ gleich ist) immer zu $\mathfrak{N}_{\bar{p}}$, also ist $SP_{\mathfrak{L}}f = 0 = P_{\mathfrak{L}}Sf$ ($P_{\mathfrak{L}}$ ist der Projektions-Operator von $\mathfrak{L}$, vgl. a. a. O.). Also reduziert auch $\mathfrak{N}_{\bar{p}} = \mathfrak{L} - \mathfrak{L} S$; die in $\mathfrak{N}_{\bar{p}}$, $\mathfrak{L}$ liegenden Teile von S seien S' bzw. S'' . Der Wertevorrat von S ist $\mathfrak{N}_{\bar{p}}$; wegen $Sf = P_{\mathfrak{N}_{\bar{p}}}Sf = SP_{\mathfrak{N}_{\bar{p}}}f$ nimmt S jeden Wert schon in $\mathfrak{N}_{\bar{p}}$ an — also hat S' auch den Wertevorrat $\mathfrak{N}_{\bar{p}}$. Als ein-eindeutiges lin. Bild des unendlichvieldimensionalen (weil überall dichten) Definitionsbereiches von R_p ist auch $\mathfrak{N}_{\bar{p}}$ unendlichvieldimensional, also ist $\mathfrak{N}_{\bar{p}}$ ein Hilbertscher Raum; darum wird A. E. Satz 41 anwendbar: S' ist hypermax. S'' ist sogar $= 0$.

Es gehöre nun f zu $\mathfrak{N}_{\bar{p}}$ und Sf zu $\mathfrak{N}_q$ ($\bar{p} \leq q$). Wir setzen $f = g + h$, $\tilde{R}_{\bar{p}}g$ sinnvoll, h von $\mathfrak{L}$. Dann ist $Sf = \tilde{R}_{\bar{p}}g$, und da f, Sf beide zu $\mathfrak{N}_{\bar{p}}$ gehören, sind sie zu h orth. Daher ist

$$|f|^2 \leq |f|^2 + |h|^2 = |f - h|^2 = |g|^2, |f| \leq |g|.$$

Da Sf zu $\mathfrak{N}_q$ gehört, ist $\tilde{R}_q^{-1}Sf = \tilde{R}_q^{-1}\tilde{R}_{\bar{p}}g = g$. Also gilt

$$|f| \leq |g| = |\tilde{R}_q^{-1}Sf| \leq \frac{1}{q} \cdot |Sf|, |Sf| \geq q \cdot |f|.$$

S' ist ein H. O. in $\mathfrak{N}_{\bar{p}}$, der alle Werte annimmt, also jeden nur einmal (A. E. Satz 41), somit hat es eine überall (in $\mathfrak{N}_{\bar{p}}$!) sinnvolle Inverse S'^{-1} . S'^{-1} ist (wie S') ein H. O. und überall sinnvoll. Wenn g zu $\mathfrak{N}_q$ gehört, so können wir auf $f = S'^{-1}g$ die Überlegungen von vorhin anwenden, es ist also

$$|S'^{-1}g| = |f| \leq \frac{1}{q} \cdot |Sf| = \frac{1}{q} \cdot |S'f| = \frac{1}{q} \cdot |g|.$$

Da $\mathfrak{N}_{\bar{p}} - \mathfrak{N}_q$ höchstens $(q - \bar{p})$ -dimensional ist, können wir endlich viele norm. orth. $\varphi_1, \dots, \varphi_k$ angeben, die es aufspannen, so daß alle zu ihnen orth. g zu $\mathfrak{N}_q$ gehören müssen. q war beliebig, wir können also sagen:

S'^{-1} ist ein überall sinnvoller H. O., derart, daß zu jedem $\varepsilon > 0$ eine endliche Zahl von norm. orth. $\varphi_1, \dots, \varphi_k$ angegeben werden kann, so daß für alle zu demselben orth. f

$$|S'^{-1}f| \leq \varepsilon \cdot |f|$$

gilt.

4. S'^{-1} ist also im Sinne der Hilbertschen Definition vollstetig, und daher gilt für dasselbe nach Hilbert: Es gibt ein vollst. norm. orth. System $\psi_1, \psi_2, \dots$ (in $\mathfrak{N}_{\bar{p}}$, d. h. ein norm. orth. System, das die abg. lin. M. $\mathfrak{N}_{\bar{p}}$ aufspannt) und eine Folge reeller Zahlen $\alpha_1, \alpha_2, \dots$, derart daß $S'^{-1}\psi_\mu = \alpha_\mu \psi_\mu$, $\mu = 1, 2, \dots$, gilt¹⁸⁾. Sei außerdem $\chi_1, \dots, \chi_s$ ein norm. orth. System, daß $\mathfrak{L}$ aufspannt (es ist eventuell $s = \infty$, aber wegen $\mathfrak{L} \neq (0)$ gewiß $s \neq 0$), dort gilt $S''\chi_\mu = 0$. Da S'^{-1} eine Inverse ist, sind alle $\alpha_\mu \neq 0$, daher ist

$S'\psi_\mu = \frac{1}{\alpha_\mu} \cdot \psi_\mu$. Wir haben also:

¹⁸⁾ Vgl. etwa Hilbert, Gött. Nachr. 1906, Seite 203, Satz VI.

$$S\psi_\mu = \frac{1}{\alpha_\mu} \cdot \psi_\mu, \quad S\chi_\mu = 0.$$

Die $\psi_1, \psi_2, \dots, \chi_1, \dots, \chi_s$ nennen wir $\varphi_1, \varphi_2, \dots$, entsprechend $\alpha_1, \alpha_2, \dots, 0, \dots, 0, a_2, \dots$. Unter den $a_1, a_2, \dots$ kommt dann (wegen $s \neq 0$) sicher die 0 vor, und es ist $S\varphi_\mu = a_\mu \varphi_\mu$. Die $\varphi_1, \varphi_2, \dots$ spannen dieselbe abg. lin. M. auf, wie $\mathfrak{R}_{\bar{p}}, \mathfrak{L}$, also $\mathfrak{L}$; daher bilden sie ein vollst. norm. orth. System. Da S', S'' abg. lin. sind, ist es S gleichfalls. Aus alledem folgt nach dem Schlußresultat des § 1 dieses Kapitels, daß S der abg. lin. H. O. von $\{a_\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$ ist.

Dabei ist R Forts. von R_p, S ebenfalls (nämlich von $\tilde{R}_p$), also sind R, S benachbart. Nun können wir beweisen:

Satz 8. Ein nicht beschränkter abg. lin. H. O. R ist immer einem Diagonaloperator S , d. h. dem abg. lin. H. O. eines $\{a_\mu \delta_{\mu\nu}\}, \{\psi_1, \psi_2, \dots\}$ benachbart. Dabei kann a_1 willkürlich vorgeschrieben werden.

Beweis: R ist nach oben oder nach unten nicht halbbeschränkt.

Im ersteren Falle, und wenn $a_1 = 0$ gewünscht wird, haben wir soeben alles erledigt; allerdings könnte für ein anderes μ $a_\mu = 0$ sein, dann genügt es aber a_1, a_μ sowie φ_1, φ_μ miteinander zu vertauschen. Alle anderen Fälle sind durch Betrachten von $R - a_1 \cdot 1, S - a_1 \cdot 1$ bzw. $-R + a_1 \cdot 1, -S + a_1 \cdot 1$ an Stelle von R, S hierauf zurückführbar.

Die mit Satz 6 einsetzende Pathologie erscheint durch diesen Satz bedeutend verschärft. Zwei k. u. Transformationen genügen (wegen der Benachbarkeit), um jede nicht beschränkte H. quadr. Matrix A auf die Diagonalform zu bringen, und dabei kann sogar ein Eigenwert willkürlich vorgeschrieben werden! (Wenn es in einer k. u. Transformation gehen sollte, so müßte, da jeder Diagonaloperator hypermax. ist und somit Satz 4 anwendbar wird, der Operator von A eine hypermax. Forts. besitzen, die Diagonaloperator ist — was natürlich eine wesentliche Aussage über deren „Spektralform“ [in der Terminologie von A. E. Satz 36, ihre Z. d. E.] involviert, nämlich daß kein „Streckenspektrum“ existiert. Sobald man aber zwei Schritte zuläßt, geht es immer!) Aber es wird bald noch schlimmer.

5. Wenn R ein willkürlicher und U ein unitärer Operator ist, so verstehen wir unter URU^{-1} denjenigen Operator, der Sinn hat, wer $\vee R(U^{-1}f)$ Sinn hat, und dann den Wert $URU^{-1}f$ annimmt. Wenn R eine der Eigenschaften hat: abg., lin., H. O., max., hypermax., definit, beschränkt, halbbeschränkt, so gilt dasselbe von URU^{-1} , analog steht es mit der Forts. (wo R, S durch URU^{-1}, USU^{-1} zu ersetzen sind). — Ferner sagen wir, daß ein nicht beschränkter H. O. R in unendlich viele Teile zerlegbar ist, wenn es eine Folge $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ unendlichvieldimensionaler abg. lin. M. mit den folgenden Eigenschaften gibt: Je zwei sind zueinander orth., sie spannen zusammen die abg. lin. M. $\mathfrak{L}$ auf, jede reduziert R , für keine ist der in ihr liegende Teil von R beschränkt.

Wir zeigen zuerst: Jeder nicht beschränkte max. H. O. R ist Forts. eines nicht beschränkten abg. lin. H. O. S , der in unendlich viele Teile zerlegbar ist (R selbst kann irreduzibel sein; vgl. A. E. Satz 39).

Sei erstens R nicht hypermax. Dann gibt es nach A. E. Satz 38 eine abg. lin. $\mathfrak{M}$ (unendlichvieldimensional!), die R reduziert, und in der ein $\bar{R}$ oder $-\bar{R}$ isomorpher Teil von R liegt (zur Definition von $\bar{R}$ vgl. A. E. Satz 37). Es genügt dann, unsere

Behauptung für $\bar{R}$ und $-\bar{R}$ zu beweisen, denn wenn die dazugehörigen abg. lin. M. $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ sind (alles in $\mathfrak{M}$!) und der fragliche Operator (dessen Forts. $\bar{R}$ bzw. $-\bar{R}$ ist) S' , so genügt es $\mathfrak{M}_1$ durch $\mathfrak{M}_1 + (\mathfrak{S} - \mathfrak{M})$ zu ersetzen, und S durch denjenigen Operator, der aus S' in $\mathfrak{M}$ und dem auf $\mathfrak{S} - \mathfrak{M}$ eingeschränkten R in $\mathfrak{S} - \mathfrak{M}$ nach Anweisung von A. E. Satz 21 entsteht — und wir sind am Ziele. Ferner genügt es, unsere Behauptung für $\bar{R}$ allein zu beweisen, denn für $-\bar{R}$ folgt sie daraus; dieses aber soll im Anhang erfolgen.

Sei zweitens R hypermax., $E(\lambda)$ seine Z. d. E. (A. E. Satz 36). Wäre $E(\lambda)$ für hinreichend große wie für hinreichend kleine λ konstant (also nach A. E. Def. 17 gleich 1 bzw. 0), etwa für $\lambda \geq |C|$, so wäre (nach A. E. Satz 36)

$$|(f, Rf)| = \left| \int_{-\infty}^{\infty} \lambda d|E(\lambda)f|^2 \right| = \left| \int_{-C}^C \lambda d|E(\lambda)f|^2 \right| \leq C \int_{-C}^C d|E(\lambda)f|^2 \leq C \cdot |f|^2,$$

d. h. R beschränkt (A. E. Satz 12, γ), entgegen der Annahme. Es gibt also eine Folge $\lambda_1, \lambda_2, \dots$ mit $\lambda_1 < \lambda_2 < \dots$ oder $\lambda_1 > \lambda_2 > \dots$ mit $\lambda_n \rightarrow \infty$ bzw. $\rightarrow -\infty$ und $E(\lambda_1) \neq E(\lambda_2)$, $E(\lambda_2) \neq E(\lambda_3)$, $\dots$. Nehmen wir etwa das erstere an, im anderen Falle schließen wir ganz analog.

$E(\lambda_{n+1}) - E(\lambda_n)$ ist ein Projektions-Operator (kurz: P. O., vgl. A. E. Def. 12), etwa $P_{\mathfrak{M}_n}$ (A. E. Satz 15); da es $\neq 0$ ist, ist $\mathfrak{M}_n \neq (0)$. Ganz wie in A. E. Kap. X (vor Satz 39) beweist man, daß $\mathfrak{M}_n$ R reduziert. Für $m < n$ ist $P_{\mathfrak{M}_m}$ Teil von $E(\lambda_{m+1})$, also von $E(\lambda_n)$, und dies ist orth. zu $1 - E(\lambda_n)$, wovon $P_{\mathfrak{M}_n}$ Teil ist — daher sind $\mathfrak{M}_m, \mathfrak{M}_n$ orth., und da es für $m < n$ gilt, gilt es immer für $m \neq n$. Ferner ist für alle f von $\mathfrak{M}_n$

$$\begin{aligned} (f, Rf) &= \int_{-\infty}^{\infty} \lambda d|E(\lambda)f|^2 \geq \int_{\lambda_n}^{\lambda_{n+1}} \lambda d|E(\lambda)f|^2 \geq \lambda_n \{ |E(\lambda_{n+1})f|^2 - |E(\lambda_n)f|^2 \} \\ &= \lambda_n \cdot |f|^2, \\ \lambda_n \cdot |f|^2 &\leq (f, Rf) \leq |f| \cdot |Rf|, \quad |Rf| \geq \lambda_n \cdot |f|. \end{aligned}$$

Wir bilden nun eine Folge von Folgen $n_1^1, n_1^2, \dots; n_2^1, n_2^2, \dots; \dots$, die zusammen genau 1, 2, $\dots$ ausmachen. $\mathfrak{M}_{n_1^1}, \mathfrak{M}_{n_2^1}, \dots$ mögen die abg. lin. M. $\mathfrak{M}_n$ aufspannen, offenbar sind die $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ paarweise orth. Man überlegt sich leicht, daß jedes $\mathfrak{M}_\mu$ R reduziert (an Hand von A. E. Satz 21), außerdem ist es klar, daß jedes $\mathfrak{M}_\mu$ unendlichvioldimensional ist. Weiter ist in $\mathfrak{M}_{n_k^\mu} \quad |Rf| \geq \lambda_{n_k^\mu} |f|$, wegen $\lambda_{n_k^\mu} \rightarrow \infty$ (für $k \rightarrow \infty$), also ist der in $\mathfrak{M}_\mu$ liegende Teil von R unbeschränkt. Unsere $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ leisten somit alles verlangte (mit $S = R|$, S mußte nur für den ersten Fall eingeführt werden, bei $\bar{R}$, vgl. den Anhang), nur daß die von ihnen aufgespannte abg. lin. M. $\mathfrak{M}$ von $\mathfrak{S}$ verschieden sein kann. Aber dann ersetzen wir $\mathfrak{M}_1$ durch $\mathfrak{M}_1 + (\mathfrak{S} - \mathfrak{M})$ (mit $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ reduziert auch $\mathfrak{M}$ R , also auch $\mathfrak{S} - \mathfrak{M}$ und $\mathfrak{M}_1 + (\mathfrak{S} - \mathfrak{M})$), und alles ist erreicht.

6. Nach diesen Vorbereitungen beweisen wir:

Satz 9. Zu zwei beliebigen nicht beschränkten max. H. O. R, S gibt es immer ein unitäres U , so daß R, USU^{-1} in zwei Schritten benachbart sind.

Beweis: Nach den Überlegungen des § 5 gilt es zwei abg. lin. H. O. R', S' , von denen R, S Forts. sind, sowie zwei Folgen unendlichvioldimensionaler abg. lin. M. $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ und $\mathfrak{N}_1, \mathfrak{N}_2, \dots$ mit den folgenden Eigenschaften: Sowohl alle $\mathfrak{M}_n$ als auch alle $\mathfrak{N}_n$

spannen die abg. lin. M. $\mathfrak{H}$ auf, sowohl $\mathfrak{M}_m, \mathfrak{M}_n$, als auch $\mathfrak{N}_m, \mathfrak{N}_n$ sind für $m \neq n$ orth.¹⁹⁾; alle $\mathfrak{M}_n$ bzw. $\mathfrak{N}_n$ reduzieren R' bzw. S' , und die in ihnen liegenden Teile R_n bzw. S_n derselben sind nicht beschränkt. Nach Satz 8 ist also jedes R_n und S_n einem Diagonaloperator Δ_n bzw. H_n (den abg. lin. H. O. eines $\{a_\mu^n \delta_{\mu\nu}\}, \{\psi_1^n, \psi_2^n, \dots\}$ bzw. $\{b_\mu^n \delta_{\mu\nu}\}, \{\chi_1^n, \chi_2^n, \dots\}$) benachbart — wobei wir noch alle a_1^n, b_1^n ($n = 1, 2, \dots$) willkürlich vorschreiben können. Wenn wir aus den $\Delta_1, \Delta_2, \dots$ bzw. aus den $H_1, H_2, \dots$ nach A. E. Satz 21 einen abg. lin. H. O. Δ bzw. H zusammensetzen, so sieht man mühelos ein, daß dieser R' bzw. S' benachbart sein wird — also auch ihren Forts. R bzw. S .

Ferner ist

$$\Delta \psi_\mu^n = a_\mu^n \psi_\mu^n, \quad H \chi_\mu^n = b_\mu^n \chi_\mu^n \quad (n, \mu = 1, 2, \dots),$$

und die ψ_μ^n sowohl als auch die χ_μ^n bilden je ein norm. orth. System; die von ihnen aufgespannte abg. lin. M. ist somit dieselbe wie die der $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ bzw. der $\mathfrak{N}_1, \mathfrak{N}_2, \dots$, also beidemal $\mathfrak{H}$. Also sind beide Systeme vollst. Wir ordnen $a_\mu^n, \psi_\mu^n, b_\mu^n, \chi_\mu^n$ zu einfachen Folgen $a_\varrho, \psi_\varrho, b_\varrho, \chi_\varrho$ um, dann sind (nach den Resultaten vom § 1 dieses Kapitels) Δ und H die abg. lin. H. O. von $\{a_\varrho \delta_{\varrho\sigma}\}, \{\psi_1, \psi_2, \dots\}$ bzw. $\{b_\varrho \delta_{\varrho\sigma}\}, \{\chi_1, \chi_2, \dots\}$. Sie sind also Diagonaloperatoren, und die den a_1^n bzw. b_1^n entsprechenden a_ϱ bzw. b_ϱ können wir nach Belieben vorschreiben. Dabei ist von der Wahl eines a_1^n bzw. b_1^n nur das entsprechende Δ_n bzw. H_n abhängig, also die a_ϱ^n bzw. b_ϱ^n ($\varrho \neq 1$) — die übrigen a_ϱ^m bzw. b_ϱ^m ($m \neq n, \varrho$ beliebig) sowie die b_ϱ^m bzw. a_ϱ^m (m, ϱ beliebig) dagegen nicht.

Wir werden nun erzwingen, daß die a_ϱ^n und b_ϱ^n , also die a_ϱ und b_ϱ , bis auf die Reihenfolge übereinstimmen. Wir teilen die Zahlenfolge $1, 2, \dots$ in eine Folge von Folgen $F_1, F_2, \dots$ ein, und verfahren dann folgendermaßen: Die a_1^n, n aus F_1 , wählen wir beliebig; dadurch werden alle a_ϱ^n, n aus F_1, ϱ beliebig, festgelegt, die wir in eine Folge ordnen. Die b_1^n, n aus F_1 wählen wir diesen a_ϱ^n gleich; dadurch werden außerdem alle b_ϱ^n, n aus $F_1, \varrho \neq 1$, festgelegt, die wir in eine Folge ordnen. Die a_1^n, n aus F_2 , wählen wir diesen b_ϱ^n gleich; dadurch werden außerdem alle a_ϱ^n, n aus $F_2, \varrho \neq 1$, festgelegt, die wir in eine Folge ordnen. Die b_1^n, n aus F_2 , wählen wir diesen a_ϱ^n gleich; dadurch werden außerdem alle b_ϱ^n, n aus $F_2, \varrho \neq 1$, festgelegt, die wir in eine Folge ordnen. Usw., usw., so werden sukzessive alle a_μ^n, b_μ^n bestimmt, und sie stimmen in der Tat bis auf die Reihenfolge überein — ebenso die a_ϱ, b_ϱ .

Durch eine Umordnung der b_ϱ und ψ_ϱ (die an H nichts ändert), können wir sogar $a_\varrho = b_\varrho$ ($\varrho = 1, 2, \dots$) erreichen. Wenn also U diejenige unitäre Transformation ist, die $\chi_1, \chi_2, \dots$ in $\psi_1, \psi_2, \dots$ überführt (daß es eine solche gibt, folgt z. B. aus A. E. Satz 26), so erkennt man: Δ bzw. H ist der abg. lin. H. O. von $\{a_\varrho \delta_{\varrho\sigma}\}, \{\psi_1, \psi_2, \dots\}$ bzw. $\{a_\varrho \delta_{\varrho\sigma}\}, \{U^{-1} \psi_1, U^{-1} \psi_2, \dots\}$, also ist $\Delta = U H U^{-1}$. Dabei sind R, Δ benachbart, ebenso S, H , und daher auch $U S U^{-1}, U H U^{-1} = \Delta$; also sind $R, U S U^{-1}$, wie behauptet wurde, in zwei Schritten benachbart.

Satz 10. Zwei beliebige nicht beschränkte H. quadr. Matrizen A, B sind immer in drei Schritten k. u. äquivalent. Die Zahl drei läßt sich im allgemeinen nicht mehr verringern.

Beweis: Sei $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System, R und S die abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ bzw. $B, \{\psi_1, \psi_2, \dots\}$. Seien R', S' je eine max. Forts. von R bzw. S . Auf R', S' ist Satz 9 anwendbar: es gibt ein unitäres U , so daß $R', U S' U^{-1}$ in zwei Schritten benachbart sind, dabei ist $U S' U^{-1}$ eine max. Forts. von $U S U^{-1}$. Wenn wir

¹⁹⁾ Natürlich nicht notwendig die $\mathfrak{M}_m$ zu den $\mathfrak{N}_n$!

Satz 5. Es ist stets $R \sim R$. Aus $R \sim S$ folgt immer $S \sim R$. Dagegen folgt bei festem S dann und nur dann für alle R, T aus $R \sim S, S \sim T$ die Relation $R \sim T$, wenn S max. ist.

Beweis: Die zwei ersten Behauptungen sind auf Grund von Satz 3 evident; wir wenden uns der dritten zu.

Sei S max., $R \sim S, S \sim T$. Dann haben R, S und T, S je eine gemeinsame Forts. S' bzw. S'' . Da S max. ist, muß $S' = S'' = S$ sein, also haben R, T die gemeinsame Forts. S , es ist $R \sim T$. — Sei umgekehrt S nicht max. Dann hat es mehrere max. Forts. (A. E. Satz 35 und Zusatz); seien R, T zwei verschiedene. R, S und T, S haben eine gemeinsame Forts. (R bzw. T), daher ist $R \sim S, S \sim T$; dagegen sind R, T max. und verschieden, also gewiß nicht $R \sim T$ (Satz 4). —

Diese Sätze zeigen, daß sich bei der k. u. Äquivalenz die beschränkten H. O., und nur diese, ganz einwandfrei benehmen. Bei max. H. O. ist auch noch einiges günstig: zwei solche sind nur k. u. äquivalent wenn sie identisch sind (allerdings kann einer einem nicht max. äquivalent sein!), für sie als Mittelglied der Kette gilt das transitive Gesetz (aber nicht, wenn sie nur Endglieder sind, wie wir bald sehen werden). Aber im allgemeinen zeigt z. B. Satz 5, daß dem beschränkten Falle völlig unähnliche Verhältnisse zu erwarten sind.

2. Auf Grund der Konvention im § 3 von Kap. I haben wir eine Matrix A als definit zu bezeichnen, wenn der abg. lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ definit ist. Man sieht sofort, daß dieser es dann und nur dann ist, wenn es der lin. H. O. von $A, \{\varphi_1, \varphi_2, \dots\}$ ist; und die Definität dieses Operators $\hat{R}$ bedeutet: es ist immer $(\hat{R}f, f) \geq 0, f = \sum_1^N x_\mu \varphi_\mu$, d. h.

$$\sum_1^N \left(\sum_1^N a_{\mu\nu} x_\mu \right) \bar{x}_\nu = \sum_1^N \sum_1^N a_{\mu\nu} x_\mu \bar{x}_\nu \geq 0.$$

Also müssen alle Abschnitte von A , das sind die endlichen Matrizen $\{a_{\mu\nu}\}, \mu, \nu = 1, \dots, N$, ($N = 1, 2, \dots$), positiv-semidefinit sein.

Eine nach oben nicht halbbeschränkte H. quadr. Matrix A hat (mit irgendeinem $\{\varphi_1, \varphi_2, \dots\}$) einen ebensolchen abg. lin. H. O. R . Nach A. E. Satz 46 ist R Forts. eines definiten lin. H. O. S (den wir durch $\tilde{S}$ ersetzen, also gleich abg. annehmen können), es ist also $R \sim S$. S möge etwa zu $B, \{\psi_1, \psi_2, \dots\}$ gehören, dann kann A k. u. in B transformiert werden — und nach dem vorhin Gesagten hat B lauter positiv-semidefinite Abschnitte.

Dieses Resultat kann noch verschärft werden:

Satz 6. Wenn A eine nach oben bzw. unten nicht halbbeschränkte H. quadr. Matrix ist, und C eine beliebige Zahl, so kann A k. u. in eine Matrix B transformiert werden, von der alle Abschnitte (d. h. die endlichen Matrizen $\{b_{\mu\nu}\}, \mu, \nu = 1, \dots, N; B = \{b_{\mu\nu}\}, \mu, \nu = 1, 2, \dots$) lauter Eigenwerte $\geq C$ bzw. $\leq C$ haben.

Beweis: Für nicht Halbbeschränktheit nach oben und $C = 0$ haben wir dies gerade vorhin festgestellt, die übrigen Fälle sind durch Betrachten von $A - C \cdot 1, B - C \cdot 1$ bzw. $-A + C \cdot 1, -B + C \cdot 1$ statt A, B darauf zurückführbar.

¹⁴⁾ 1 sei die Einheitsmatrix $\{\delta_{\mu\nu}\}, \delta_{\mu\nu} = \begin{cases} 1 & \text{für } \mu = \nu, \\ 0 & \text{für } \mu \neq \nu. \end{cases}$

$N_1, N_2, \dots$ von $1, 2, \dots$, sowie zwei andere elementfremde Teilfolgen $M'_1, M'_2, \dots$ und $N'_1, N'_2, \dots$ von $1, 2, \dots$, so daß

$$a_{M_n} = 0, \quad a_{N_n} \rightarrow \infty, \quad b_{M'_n} \rightarrow \infty, \quad b_{N'_n} = 0, \quad \varphi_{M_n} = \varphi_{M'_n}, \quad \varphi_{N_n} = \varphi_{N'_n} \quad (n = 1, 2, \dots)$$

gilt. Dann sind R, S in zwei Schritten benachbart.

Beweis: Seien $\mathfrak{C}, \mathfrak{D}$ zwei elementfremde, überall dichte Teilmengen der Definitionsbereiche von R bzw. S ; T sei derjenige Operator, der in $\mathfrak{C}$ mit R und in $\mathfrak{D}$ mit S übereinstimmend definiert ist. Wenn T der Bedingung

$$(Tf, g) = (f, Tg)$$

genügt, so ist es ein H. O., also $\tilde{T}$ ein abg. lin. H. O. $\tilde{T}$ ist dann R und S benachbart, also sind R, S in zwei Schritten benachbart, wie behauptet wurde. Es gilt also nur, die obige Relation zu sichern.

Wenn f, g beide zu $\mathfrak{C}$ oder beide zu $\mathfrak{D}$ gehören, so gilt sie gewiß, es bleibt also nur der Fall zu erörtern, wo eines zu $\mathfrak{C}$ und das andere zu $\mathfrak{D}$ gehört. D. h. daß für alle f von $\mathfrak{C}$ und g von $\mathfrak{D}$

$$(Rf, g) = (f, Sg)$$

gilt. Wir werden $\mathfrak{C}$ und $\mathfrak{D}$ demgemäß bestimmen, u. zw. als überall dichte Folgen $f_1, f_2, \dots$ und $g_1, g_2, \dots$.

Wir ordnen alle $\sum_1^n a_\sigma \varphi_\sigma$ sowie alle $\sum_1^n a_\sigma \psi_\sigma$ ($n = 1, 2, \dots$, die $a_1, \dots, a_n$ rational)

in je eine Folge $\bar{f}_1, \bar{f}_2, \dots$ bzw. $\bar{g}_1, \bar{g}_2, \dots$; beide Folgen sind offenbar überall dicht, und alle $R\bar{f}_\mu$ bzw. $S\bar{g}_\nu$ sind sinnvoll. Da wir die M_n, M'_n und N_n, N'_n durch irgendwelche Teilfolgen M_{p_n}, M'_{p_n} und N_{p_n}, N'_{p_n} ersetzen können (unbeschadet der im Satze formulierten Eigenschaften), also insbesondere durch solche mit $|b_{M'_{p_n}}| \geq C_n, |a_{N_{p_n}}| \geq C_n$ — so dürfen wir annehmen, daß von vornherein $|b_{M'_n}| \geq C_n, |a_{N_n}| \geq C_n$ galt. Über die $C_1, C_2, \dots$ werden wir noch verfügen, natürlich unabhängig von den M_n, M'_n, N_n, N'_n und den nun weiter zu entwickelnden Begriffen, die von diesen abhängen.

An Stelle der einfachen Indizierung $C_n, M_n, M'_n, N_n, N'_n$ ($n = 1, 2, \dots$) führen wir eine dreifache ein: $C_{\mu\nu}^e, M_{\mu\nu}^e, M'_{\mu\nu}^e, N_{\mu\nu}^e, N'_{\mu\nu}^e$ ($\mu, \nu, e = 1, 2, \dots$)²⁴). Das größte e , für welches in $\bar{f}_\mu = \sum_1^n a_\sigma \varphi_\sigma$ ein $a_{M_{\mu\nu}^e} \neq 0$ vorkommt, heiße $F(\mu)$; das größte e , für welches in $\bar{g}_\nu = \sum_1^n a_\sigma \psi_\sigma$ ein $a_{N_{\mu\nu}^e} \neq 0$ vorkommt, heiße $G(\nu)$. Ferner soll $C_{\mu\nu}^e$ nur von e abhängen: $C_{\mu\nu}^e = D_e$. Wir setzen nun:

$$f_\mu = \bar{f}_\mu + \sum_1^{\infty} \frac{1}{\mu\nu e} \varphi_{M_{\mu\nu}^e} + \sum_1^{\mu} x_{\mu\nu} \varphi_{N_{\mu\nu}^e(\mu, \nu)},$$

$$g_\nu = \bar{g}_\nu + \sum_1^{\infty} \frac{1}{\mu\nu e} \psi_{N_{\mu\nu}^e} + \sum_1^{\nu-1} x_{\mu\nu} \psi_{M_{\mu\nu}^e(\mu, \nu)}.$$

($\sum_1^{\infty} \frac{1}{\mu\nu e}, \sum_1^{\mu} x_{\mu\nu}$ sind offenbar konvergent). Über die $x_{\mu\nu}$ und $e(\mu, \nu)$ (die für $\mu \geq \nu$ in f_μ und für $\mu < \nu$ in g_ν eingehen) werden wir noch verfügen, jedenfalls sei aber

²⁴) Die Triplets μ, ν, e ($= 1, 2, \dots$) sollen also den n ($= 1, 2, \dots$) ein-eindeutig zugeordnet werden.

$$|x_{\mu\nu}| \leq \frac{1}{\mu\nu}.$$

Zunächst ist (wir setzen das endliche $\sum_j \frac{1}{\sigma^2} = s$)

$$|f_\mu - \bar{f}_\mu|^2 = \sum_1^\infty e \frac{1}{\mu^2 \nu^2 \varrho^2} + \sum_1^\mu |x_{\mu\nu}|^2 \leq \sum_1^\infty e \frac{1}{\mu^2 \nu^2 \varrho^2} + \sum_1^\infty \frac{1}{\mu^2 \nu^2} = \frac{s^2 + s}{\mu^2},$$

$$|g_\nu - \bar{g}_\nu|^2 = \sum_1^\infty e \frac{1}{\mu^2 \nu^2 \varrho^2} + \sum_1^{\nu-1} |x_{\mu\nu}|^2 \leq \sum_1^\infty e \frac{1}{\mu^2 \nu^2 \varrho^2} + \sum_1^\infty \frac{1}{\mu^2 \nu^2} = \frac{s^2 + s}{\nu^2},$$

also ist $f_\mu - \bar{f}_\mu \rightarrow 0$, $g_\nu - \bar{g}_\nu \rightarrow 0$ (für $\mu \rightarrow \infty$ bzw. $\nu \rightarrow \infty$). Hieraus, und daraus daß die $\bar{f}_1, \bar{f}_2, \dots$ und die $\bar{g}_1, \bar{g}_2, \dots$ überall dicht liegen, folgt, daß auch die $f_1, f_2, \dots$ und die $g_1, g_2, \dots$ überall dicht sind. Ferner haben offenbar die Rf_μ und die Sg_ν Sinn:

$$Rf_\mu = R\bar{f}_\mu + \sum_1^\mu a_{N_{\mu\nu}^e(\mu, \nu)} x_{\mu\nu} \cdot \varphi_{N_{\mu\nu}^e(\mu, \nu)}, \quad Sg_\nu = S\bar{g}_\nu + \sum_1^{\nu-1} b_{M_{\mu\nu}^e(\mu, \nu)} x_{\mu\nu} \cdot \varphi_{M_{\mu\nu}^e(\mu, \nu)}.$$

Wir können darum, wenn wir noch $\varphi_{M_{\mu\nu}^e} = \varphi_{M_{\mu\nu}^e}$, $\varphi_{N_{\mu\nu}^e} = \varphi_{N_{\mu\nu}^e}$ berücksichtigen, die Gleichung $(Rf_\mu, g_\nu) = (f_\mu, Sg_\nu)$ hinschreiben. Und zwar lautet sie für $\mu \geq \nu$

$$\begin{aligned} & (R\bar{f}_\mu, \bar{g}_\nu + \sum_1^\infty e \frac{1}{x\nu\varrho} \varphi_{N_{x\nu}^e}) + \sum_1^{\nu-1} \bar{x}_{x\nu} (R\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}) + \sum_1^\mu a_{N_{\mu\lambda}^e(\mu, \lambda)} x_{\mu\lambda} (\varphi_{N_{\mu\lambda}^e(\mu, \lambda)}, \bar{g}_\nu) \\ & + \frac{1}{\mu\nu \cdot \varrho(\mu, \nu)} a_{N_{\mu\nu}^e(\mu, \nu)} x_{\mu\nu} \\ & = (\bar{f}_\mu + \sum_1^\infty e \frac{1}{\mu\lambda\varrho} \varphi_{N_{\mu\lambda}^e}, S\bar{g}_\nu) + \sum_1^\mu x_{\mu\lambda} (\varphi_{N_{\mu\lambda}^e(\mu, \lambda)}, S\bar{g}_\nu) + \sum_1^{\nu-1} b_{M_{x\nu}^e(x, \nu)} x_{x\nu} (\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}), \end{aligned}$$

und für $\mu < \nu$

$$\begin{aligned} & (R\bar{f}_\mu, \bar{g}_\nu + \sum_1^\infty e \frac{1}{n\nu\varrho} \varphi_{N_{x\nu}^e}) + \sum_1^{\nu-1} \bar{x}_{x\nu} (R\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}) + \sum_1^\mu a_{N_{\mu\lambda}^e(\mu, \lambda)} x_{\mu\lambda} (\varphi_{N_{\mu\lambda}^e(\mu, \lambda)}, \bar{g}_\nu) \\ & = (\bar{f}_\mu + \sum_1^\infty e \frac{1}{\mu\lambda\varrho} \varphi_{N_{\mu\lambda}^e}, S\bar{g}_\nu) + \sum_1^\mu x_{\mu\lambda} (\varphi_{N_{\mu\lambda}^e(\mu, \lambda)}, S\bar{g}_\nu) + \sum_1^{\nu-1} b_{M_{x\nu}^e(x, \nu)} x_{x\nu} (\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}) \\ & + \frac{1}{\mu\nu \cdot \varrho(\mu, \nu)} b_{M_{\mu\nu}^e(\mu, \nu)} x_{\mu\nu}. \end{aligned}$$

Wenn die allgemeine Vorschrift

$$\varrho(x, \lambda) > \begin{cases} \text{Max } (G(1), \dots, G(\lambda)) \\ \text{Max } (F(1), \dots, F(x)) \end{cases}$$

aufgestellt wird, so dürfen wir die in unseren Formeln vorkommenden vier $\sum_1^\mu$ durch $\sum_1^{\nu-1}$ ersetzen, und die vier $\sum_1^{\nu-1}$ durch $\sum_1^{\mu-1}$, wenigstens soweit das erste Summationszeichen weiter reicht als das zweite, d. h. für $\mu \geq \nu - 1$ bzw. $\mu \leq \nu$ — also allenfalls in der ersten bzw. zweiten Formel. (Denn nach Definition der $F(x)$, $G(\lambda)$ haben wir für $x \geq \mu$ bzw. $\lambda \geq \nu$

$$(\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}) = 0,$$

$$(R\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}) = (\bar{f}_\mu, R\varphi_{M_{x\nu}^e(x, \nu)}) = a_{x\nu}^e(x, \nu) (\bar{f}_\mu, \varphi_{M_{x\nu}^e(x, \nu)}) = 0,$$

und ebenso

$$(\varphi_{N_{\mu\lambda}^e(\mu, \lambda)}, \bar{g}_\nu) = (\varphi_{N_{\mu\lambda}^e(\mu, \lambda)}, S\bar{g}_\nu) = 0$$

wegen $\varrho(\kappa, \nu) > F(\mu)$ bzw. $\varrho(\mu, \lambda) > G(\nu)$, wobei wieder $\varphi_{N^{\varrho(\mu, \lambda)}} = \varphi_{N^{\varrho(\mu, \lambda)}}$ zu beachten ist.) Infolgedessen kommen in der ersten Formel außer $\varrho(\mu, \nu)$, $x_{\mu\nu}$ nur solche $\varrho(\kappa, \lambda)$, $x_{\kappa\lambda}$ vor, für die $\kappa < \nu \leq \mu$, $\lambda = \nu$ oder $\kappa = \mu$, $\lambda < \nu$ ist, also immer $\kappa + \lambda < \mu + \nu$; und in der zweiten Formel ebenfalls.

Wir werden darum die $\varrho(\mu, \nu)$, $x_{\mu\nu}$ in der Reihenfolge wachsender $\mu + \nu$ bestimmen (für $\mu \geq \nu$ ist immer die erste Gleichung zu berücksichtigen, für $\mu < \nu$ die zweite) — dann können wir bei jedem Schritt über $\varrho(\mu, \nu)$ frei verfügen, und $x_{\mu\nu}$ ist durch $\varrho(\mu, \nu)$ und die bereits erledigten $\varrho(\kappa, \lambda)$, $x_{\kappa\lambda}$ ($\kappa + \lambda < \mu + \nu$) eindeutig festgelegt. Wir brauchen nur noch die Bedingung $|x_{\mu\nu}| \leq \frac{1}{\mu\nu}$ durch geeignete Wahl von $\varrho(\mu, \nu)$ zu sichern.

Die Gleichung für $x_{\mu\nu}$ hat die Gestalt

$$\frac{1}{\mu\nu \cdot \varrho(\mu, \nu)} a_{N^{\varrho(\mu, \nu)}} x_{\mu\nu} = A \text{ bzw. } \frac{1}{\mu\nu \cdot \varrho(\mu, \nu)} b_{M^{\varrho(\mu, \nu)}} x_{\mu\nu} = A,$$

wobei A nur von den $\varrho(\kappa, \lambda)$, $x_{\kappa\lambda}$ mit $\kappa + \lambda < \mu + \nu$ abhängt, nicht aber von $\varrho(\mu, \nu)$, $x_{\mu\nu}$. Daher ist

$$|x_{\mu\nu}| = \frac{\mu\nu \cdot |A| \cdot \varrho(\mu, \nu)}{|a_{N^{\varrho(\mu, \nu)}}|} \leq \mu\nu \cdot |A| \cdot \frac{\varrho(\mu, \nu)}{D_{\varrho(\mu, \nu)}} \text{ bzw. } |x_{\mu\nu}| = \frac{\mu\nu \cdot |A| \cdot \varrho(\mu, \nu)}{|b_{M^{\varrho(\mu, \nu)}}|} \leq \\ \leq \mu\nu \cdot |A| \cdot \frac{\varrho(\mu, \nu)}{D_{\varrho(\mu, \nu)}}.$$

Wir müssen also $\varrho(\mu, \nu) = \varrho$ so wählen, daß $\frac{D_{\varrho}}{\varrho} \geq \mu^2 \nu^2 \cdot |A|$ wird, was ohne weiteres geht, wenn $\frac{D_{\varrho}}{\varrho} \rightarrow \infty$, also wenn wir z. B. $D_{\varrho} = \varrho^2$ setzen. (Hierbei sind unendlich viele Werte von ϱ verfügbar, wodurch auch noch $f_{\mu} \neq g_{\nu}$ gewährleistet werden kann — d. h. die am Anfang erwähnte Elementfremdheit von $\mathfrak{E}, \mathfrak{D}$.)

Damit ist alles bewiesen.

2. Ehe wir aus Satz 11 weitergehende Konsequenzen ziehen, müssen wir eine Hilfsüberlegung über norm. orth. Systeme anstellen.

Satz 12. Seien k norm. orth. Systeme $\varphi_1^x, \varphi_2^x, \dots$ gegeben ($x = 1, \dots, k$, zueinander brauchen sie nicht orth. zu sein); diese besitzen immer k Teilfolgen $\varphi_{N_1}^x, \varphi_{N_2}^x, \dots$ ($x = 1, \dots, k$), derart, daß, wenn die $\varphi_{N_n}^x$ ($n = 1, \dots, k$, $n = 1, 2, \dots$) die abg. lin. M. $\mathfrak{M}$ aufspannen, $\mathfrak{S} - \mathfrak{M}$ unendlichvieldimensional ist.

Beweis: Sei $\varphi_1^{k+1}, \varphi_2^{k+1}, \dots$ ein weiteres (willkürlich wählbares) norm. orth. System. Wir werden $k+1$ Teilfolgen $\varphi_{N_1}^x, \varphi_{N_2}^x, \dots$ ($x = 1, \dots, k+1$) angeben, so daß identisch

$$\left| \sum_{\mathfrak{I}}^N \sum_{\mathfrak{I}}^{k+1} a_{m\mathfrak{x}} \varphi_{N_m}^x \right|^2 \geq \frac{1}{2} \sum_{\mathfrak{I}}^N \sum_{\mathfrak{I}}^{k+1} |a_{m\mathfrak{x}}|^2$$

gilt, und wollen zunächst zeigen, daß hieraus die Unendlichvieldimensionalität von $\mathfrak{S} - \mathfrak{M}$ folgt.

Nehmen wir an, es wäre p -dimensional (p endlich). Wir bilden:

$$\varphi_{N_n}^{k+1} = f_n + g_n, \quad f_n \text{ von } \mathfrak{M}, \quad g_n \text{ von } \mathfrak{S} - \mathfrak{M} \quad (n = 1, \dots, p+1).$$

Sei $\mathfrak{M}_x$ die von den $\varphi_{N_1}^x, \varphi_{N_2}^x, \dots$ aufgespannte abg. lin. M. ($x = 1, \dots, k$), dann ist $\mathfrak{M} = \mathfrak{M}_1 + \dots + \mathfrak{M}_x$, also

$$f_n = f_n^1 + \cdots + f_n^k, \quad f_n^1 \text{ von } \mathfrak{M}_1, \dots, f_n^k \text{ von } \mathfrak{M}_k \quad (n = 1, \dots, p+1).$$

Da $\mathfrak{S} - \mathfrak{M}$ p -dimensional ist, sind die $g_1, \dots, g_{p+1}$ nicht lin. unabhängig, es gibt also ein System von Zahlen $a_1, \dots, a_{p+1}$ (nicht alle = 0!) mit

$$a_1 g_1 + \cdots + a_{p+1} g_{p+1} = 0.$$

Daher ist

$$\sum_1^{p+1} a_n \varphi_{Nn}^{k+1} = \sum_1^{p+1} a_n f_n^1 + \cdots + \sum_1^{p+1} a_n f_n^k = \bar{f}^1 + \cdots + \bar{f}^k, \quad \bar{f}^1 \text{ von } \mathfrak{M}_1, \dots, \bar{f}^k \text{ von } \mathfrak{M}_k.$$

Da die $\varphi_{N1}^x, \varphi_{N2}^x, \dots$ ein norm. orth. System bilden, das die abg. lin. M. $\mathfrak{M}_x$ aufspannt, so muß $\bar{f}^x = \sum_1^\infty x_{tx} \varphi_{Nt}^x$ sein (vgl. etwa die Überlegungen in A. E., beim Beweise von Satz 13), also haben wir

$$\begin{aligned} \sum_1^N \sum_1^k x_{tx} \varphi_{Nt}^x - \sum_1^{p+1} a_n \varphi_{Nn}^{k+1} &\rightarrow \sum_1^k \bar{f}^x - \sum_1^{p+1} a_n \varphi_{Nn}^{k+1} = 0, \\ \left| \sum_1^N \sum_1^k x_{tx} \varphi_{Nt}^x - \sum_1^{p+1} a_n \varphi_{Nn}^{k+1} \right|^2 &\rightarrow 0 \quad (\text{für } N \rightarrow \infty). \end{aligned}$$

Aus der vorausgesetzten Ungleichung folgt aber

$$\left| \sum_1^N \sum_1^k x_{tx} \varphi_{Nt}^x - \sum_1^{p+1} a_n \varphi_{Nn}^{k+1} \right|^2 \geq \frac{1}{2} \left[\sum_1^N \sum_1^k |x_{tx}|^2 + \sum_1^{p+1} |a_n|^2 \right] \geq \frac{1}{2} \sum_1^{p+1} |a_n|^2,$$

also müßte $\sum_1^{p+1} |a_n|^2 = 0$, $a_1 = \cdots = a_{p+1} = 0$ sein, entgegen der Annahme. Es gilt nun die Gültigkeit der genannten Ungleichung zu sichern.

Es ist

$$\begin{aligned} \left| \sum_1^N \sum_1^{k+1} a_{mx} \varphi_{Nn}^x \right|^2 &= \sum_{m,n}^N \sum_1^{k+1} a_{mx} \overline{a_{n\lambda}} (\varphi_{Nn}^x, \varphi_{Nn}^x) \\ &= \sum_1^N \sum_1^{k+1} |a_{mx}|^2 + \sum_{m,n}^N \sum_{\substack{\lambda \\ (m+n)}}^{k+1} a_{mx} \overline{a_{n\lambda}} (\varphi_{Nn}^x, \varphi_{Nn}^x) \\ &\geq \sum_1^N \sum_1^{k+1} |a_{mx}|^2 - 2 \sum_{\substack{m,n \\ (m < n)}}^N |a_{mx}| |a_{n\lambda}| \cdot |(\varphi_{Nn}^x, \varphi_{Nn}^x)|. \end{aligned}$$

Für den Subtrahenden machen wir (wie schon in A. E., Beweis von Satz 46) die Abschätzung

$$\begin{aligned} &\sum_{\substack{m,n \\ (m < n)}}^N \sum_1^{k+1} |a_{mx}| |a_{n\lambda}| \cdot |(\varphi_{Nn}^x, \varphi_{Nn}^x)| \\ &\leq \sum_{\substack{m,n \\ (m < n)}}^N \sum_1^{k+1} \left(\frac{1}{8(k+1)} \frac{1}{2^{n-m}} \cdot |a_{mx}|^2 + 4(k+1) 2^{n-m} |(\varphi_{Nn}^x, \varphi_{Nn}^x)|^2 \cdot |a_{n\lambda}|^2 \right) \\ &= \sum_1^N \sum_1^{k+1} \left[\sum_1^{m-1} \sum_1^{k+1} 4(k+1) 2^{m-n} |(\varphi_{Nn}^x, \varphi_{Nn}^x)|^2 + \sum_{m+1}^N \sum_1^{k+1} \frac{1}{8(k+1)} \frac{1}{2^{n-m}} \right] \cdot |a_{mx}|^2 \\ &\leq \sum_1^N \sum_1^{k+1} \left[4(k+1) \sum_1^{m-1} \sum_1^{k+1} 2^{m-n} |(\varphi_{Nn}^x, \varphi_{Nn}^x)|^2 + \frac{1}{8} \right] \cdot |a_{mx}|^2. \end{aligned}$$

Wenn wir also dafür sorgen, daß $4(k+1) \sum_1^{m-1} \sum_1^{k+1} 2^{m-n} |(\varphi_{Nm}^\kappa, \varphi_{Nn}^\lambda)|^2 \leq \frac{1}{8}$ ist, so ist der obige Ausdruck $\leq \frac{1}{4} \sum_1^N \sum_1^{k+1} |a_{m\kappa}|^2$, also

$$\left| \sum_1^N \sum_1^{k+1} a_{m\kappa} \varphi_{Nm}^\kappa \right|^2 \leq \frac{1}{2} \sum_1^N \sum_1^{k+1} |a_{m\kappa}|^2,$$

wie wir es gewünscht hatten.

Die erwähnte Ungleichung gilt sicher, wenn etwa

$$|(\varphi_{Nm}^\kappa, \varphi_{Nn}^\lambda)|^2 \leq \frac{1}{4} \cdot \frac{1}{8(k+1)} \cdot \frac{1}{2^{m-n}} \cdot \frac{1}{2^n} = \frac{1}{(k+1)2^{m+5}}$$

ist (für $n < m$, $n, \kappa, \lambda = 1, \dots, k+1$). Dies ist aber durch geeignete Wahl der Folge $N_1, N_2, \dots$ erreichbar. Um dieselbe demgemäß zu konstruieren, müssen wir nur dieses leisten: wenn $N_1, \dots, N_{l-1}$ bereits so gewählt sind, daß die genannte Bedingung für alle $n < m \leq l-1$ gilt, N_l so bestimmen, daß sie auch für die $n < m = l$ gültig wird.

Für jedes f ist $\sum_1^\infty |(\varphi_t^\kappa, f)|^2$ endlich (A. E. Satz 4), also $|(\varphi_t^\kappa, f)|^2 \rightarrow 0$ (für $t \rightarrow \infty$).

Also ist insbesondere $|(\varphi_t^\kappa, \varphi_{Nn}^\lambda)|^2 \rightarrow 0$ ($t \rightarrow \infty$) für alle $n < l$, $\kappa, \lambda = 1, \dots, k+1$, und

darum können wir $N_l = t$ in der Tat so wählen, daß $|(\varphi_t^\kappa, \varphi_{Nn}^\lambda)|^2 \leq \frac{1}{(k+1)2^{m+5}}$ ist,

für alle $n < l$, $\kappa, \lambda = 1, \dots, k+1$. Damit ist der Beweis restlos geführt.

3. Jetzt können wir die Prämissen von Satz 11 bedeutend abschwächen.

Satz 13. Seien R, S die abg. lin. H. O. von $\{a_\mu \delta_{\mu\nu}\}$, $\{\varphi_1, \varphi_2, \dots\}$ und $\{b_\mu \delta_{\mu\nu}\}$, $\{\psi_1, \psi_2, \dots\}$. Sowohl unter den a_μ als auch unter den b_μ mögen unendlich viele Nullen vorkommen, und beide Zahlenreihen seien unbeschränkt. Dann sind R, S in sechs Schritten benachbart.

Beweis. Wir wählen zwei elementfremde Teilfolgen $p_1, p_2, \dots$ und $q_1, q_2, \dots$ aus $1, 2, \dots$ aus, mit $a_{p_n} = 0, a_{q_n} \rightarrow \infty$; und zwei andere $r_1, r_2, \dots$ und $s_1, s_2, \dots$ mit $b_{r_n} = 0, b_{s_n} \rightarrow \infty$. Wir können hier die p_n, q_n, r_n, s_n (unter Wahrung dieser Eigenschaften) durch vier beliebige Teilfolgen von sich ersetzen, nach Satz 12 können wir hierdurch erreichen, daß, wenn $\mathfrak{M}$ die von den $\varphi_{p_n}, \varphi_{q_n}, \varphi_{r_n}, \varphi_{s_n}$ aufgespannte abg. lin. M. ist, $\mathfrak{H} - \mathfrak{M}$ unendlichvieldimensional wird. Also dürfen wir annehmen, daß dies von Anfang an der Fall war.

Sei $\chi_1, \chi_2, \dots$ ein die abg. lin. M. $\mathfrak{H} - \mathfrak{M}$ aufspannendes norm. orth. System; da $\mathfrak{H} - \mathfrak{M}$ unendlichvieldimensional ist, bricht diese Folge nicht ab. Nach Konstruktion sind $\chi_1, \chi_2, \dots$ zu allen $\varphi_{p_n}, \varphi_{q_n}, \varphi_{r_n}, \varphi_{s_n}$ orth. Daher bilden die $\chi_n, \varphi_{p_n}, \varphi_{q_n}$ zusammen ein norm. orth. System, ebenso die $\chi_n, \varphi_{r_n}, \varphi_{s_n}$ — wir können beide zu je einem vollst. norm. orth. System ergänzen: $\chi_n, \varphi_{p_n}, \varphi_{q_n}, \omega'_n$ bzw. $\chi_n, \varphi_{r_n}, \varphi_{s_n}, \omega''_n$ (wobei die Anzahl der ω'_n bzw. ω''_n nicht unendlich sein muß!).

Wir numerieren die $\chi_n, \varphi_{p_n}, \varphi_{q_n}, \omega'_n$ und die $\chi_n, \varphi_{r_n}, \varphi_{s_n}, \omega''_n$ irgendwie durch: $\xi_1, \xi_2, \dots$ bzw. $\eta_1, \eta_2, \dots$. Den $\chi_n, \varphi_{p_n}, \varphi_{q_n}$ bzw. $\chi_n, \varphi_{r_n}, \varphi_{s_n}$ mögen dabei die Teilfolgen u'_n, v'_n, w'_n bzw. u''_n, v''_n, w''_n von $1, 2, \dots$ entsprechen. Schließlich bilden wir zwei Zahlenfolgen $c_1, c_2, \dots$ und $d_1, d_2, \dots$ für die jedenfalls

$$c_{v'_n} \rightarrow \infty, c_{w'_n} = 0, c_{u'_{2n-1}} = 0, c_{u'_{2n}} \rightarrow \infty; d_{v''_n} \rightarrow \infty, d_{w''_n} = 0, d_{u''_{2n-1}} \rightarrow \infty, d_{u''_{2n}} = 0$$

gelten soll. T' und T'' seien die abg. lin. H. O. von $\{c_\mu \delta_{\mu\nu}\}, \{\xi_1, \xi_2, \dots\}$ bzw. $\{d_\mu \delta_{\mu\nu}\}, \{\eta_1, \eta_2, \dots\}$.

Nun erfüllen R, T' sowie T', T'' und T'', S die Prämissen von Satz 11, und zwar müssen wir der Reihe nach setzen:

$M_n = p_n, N_n = q_n, M'_n = v'_n, N'_n = w'_n$; bzw. $M_n = u'_{2n-1}, N_n = u'_{2n}, M'_n = u'_{2n-1}, N'_n = u'_{2n}$;

bzw. $M_n = w'_n, N_n = v'_n, M'_n = s_n, N'_n = r_n$.

Daher sind sie nach dem genannten Satz in je zwei Schritten benachbart, und R, S sind es in sechs Schritten.

Satz 14. Zwei beliebige nicht beschränkte abg. lin. H. O. R, S sind immer in neun Schritten k. u. äquivalent. Diese Zahl läßt sich im allgemeinen keinesfalls unter drei verringern.

Beweis: Seien R', S' max. Forts. von R bzw. S . Nach den Überlegungen beim Beweise von Satz 9 sind R', S' gewissen Diagonaloperatoren Δ, H benachbart; es sei etwa Δ der abg. lin. H. O. von $\{a_\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$, H der von $\{b_\mu \delta_{\mu\nu}\}, \{\psi_1, \psi_2, \dots\}$. Und dabei können wir sowohl bei den a_μ als auch bei den b_μ eine ganze unendliche Teilfolge nach Belieben vorschreiben. Darum können wir bei den a_μ sowohl wie bei den b_μ erreichen, daß unendlich viele Nullen unter ihnen vorkommen, und daß die ganze Zahlenreihe unbeschränkt wird.

Nach Satz 13 sind folglich Δ, H in sechs Schritten benachbart, und somit R', S' in acht Schritten. Dabei sind R', S' max. Forts. von R bzw. S . Wie beim Beweise von Satz 10 (vgl. wie dort mit den §§ 3, 4 von Kap. II) folgern wir hieraus, daß R, S in neun Schritten k. u. äquivalent sind.

Daß die Zahl der notwendigen Schritte im allgemeinen nicht unter drei herabgesetzt werden kann, haben wir bereits am Anfang des § 1 dieses Kapitels festgestellt.

4. Auf Grund des nunmehr vorliegenden Tatsachenmaterials können wir uns von den k. u. Äquivalenz-Verhältnissen bei abg. lin. H. O. das folgende Bild machen. Wir teilen die abg. lin. H. O. in Klassen ein, so daß jede Klasse alle miteinander in endlich vielen Schritten k. u. äquivalenten H. O. enthält. Dann bilden alle unbeschränkten Operatoren eine einzige Klasse. Da aber ein beschränkter Operator nur sich selbst k. u. äquivalent ist (Satz 4), also auch in endlich vielen Schritten nur sich selbst, so bildet jeder derselben für sich allein eine ganze Klasse. Analog wie im § 7 im Kap. III sehen wir: aus der k. u. Äquivalenz in endlich vielen Schnitten folgt die in neun Schritten, aber sicher nicht immer die in weniger als drei Schritten.

Es ist schwer, dies anders zu resümieren, als durch die Feststellung: Die unitäre Transformationstheorie der Hermiteschen Matrizen löst sich, sobald man auch nur einen einzigen Schritt über den Hilbertschen Rahmen der Beschränktheit tut, völlig von der Operatorentheorie los, sie vermag deren Probleme in keiner Weise mehr adäquat wiederzugeben. Man vergegenwärtige sich etwa die Betrachtungen, die wir an Satz 6, 8, 10 sowie den gegenwärtigen Satz 14 angeschlossen haben. Dabei ist es nicht einmal unbedingt richtig, von einer Pathologie zu reden; denn die für die unbeschränkten Matrizen geltenden Sätze sind einfach und allgemein — nur lauten sie ganz anders, als man vernünftigerweise erwarten sollte. Noch eher sind diejenigen unitären Matrizen als pathologisch zu bezeichnen, die unsere paradoxen k. u. Äquivalenzen vermitteln, trotzdem gerade diese beschränkt sind! (Vgl. das Ende der Einleitung.)

V. Betrachtungen über mehrere Operatoren.

1. Bei überall sinnvollen Operatoren ist es ohne weiteres möglich, die Addition, Subtraktion und Multiplikation zu definieren:

$$(R \pm S)f = Rf \pm Sf, \quad RSf = R(Sf),$$

wovon wir ja auch häufig Gebrauch machten. Anders ist es bei nicht überall sinnvollen, also z. B. bei allen nicht beschränkten H. O., da hier die Frage des Definitionsbereiches auftritt: Um z. B. aus den H. O. R, S einen H. O. $R + S$ bilden zu können, müßte man eine überall dichte Menge von f haben, für welche Rf, Sf gleichzeitig sinnvoll sind. Im folgenden soll gezeigt werden, daß in dieser Richtung die größte Vorsicht geboten ist, weil es max. Operatoren gibt, deren Definitionsbereiche (außer der 0) kein einziges gemeinsames Element haben ²⁵⁾.

Wir zeigen zuerst mit Hilfe von Satz 11:

Satz 15. Es gibt zwei max. H. O., die für kein einziges $f \neq 0$ beide Sinn haben.

Beweis: Sei $\varphi_1, \varphi_2, \dots$ ein vollst. norm. orth. System, $a_{2n-1} = b_{2n} = 0$, $a_{2n} = b_{2n-1} = n$, R, S die abg. lin. H. O. von $\{a_\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$ bzw. $\{b_\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$. Nach Satz 11 sind sie in zwei Schritten benachbart ($M_n = 2n - 1$, $N'_n = 2n$, $M'_n = 2n - 1$, $N'_n = 2n$), seien also beide T benachbart. Wir können T als max. annehmen (§ 4, Kap. II).

Wenn Rf, Sf, Tf gleichzeitig Sinn haben, so ist wegen der Benachbarkeit von R, T und S, T $Rf = Sf = Tf$ (Satz 7). Sei $f = \sum_1^\infty x_\mu \cdot \varphi_\mu$, dann ist (§ 1, Kap. III)

$Rf = \sum_1^\infty a_\mu x_\mu \cdot \varphi_\mu$, $Sf = \sum_1^\infty b_\mu x_\mu \cdot \varphi_\mu$, also $\sum_1^\infty (a_\mu - b_\mu) x_\mu \cdot \varphi_\mu = 0$. Da stets $a_\mu \neq b_\mu$ ist, muß $x_\mu = 0$ sein, also $f = 0$.

Sei nun R' der abg. lin. H. O. von $\{\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$; als Diagonaloperator ist er hypermax. (§ 1, Kap. III, ebenso das Folgende). Wenn $R'f, f = \sum_1^\infty x_\mu \varphi_\mu$, Sinn

hat, so ist $\sum_1^\infty \mu^2 |x_\mu|^2$ endlich, also auch $\sum_1^\infty a_\mu^2 |x_\mu|^2$, $\sum_1^\infty b_\mu^2 |x_\mu|^2$ — d. h. Rf, Sf sind sinnvoll. Wenn also $R'f, Tf$ gleichzeitig Sinn haben, so muß $f = 0$ sein. Somit sind R', T die gewünschten max. H. O.

Satz 16. Es gibt einen Diagonaloperator R und einen unitären U , so daß für kein $f \neq 0$ R und URU^{-1} beide Sinn haben.

Beweis: Sei zunächst S irgendein abg. lin. H. O., etwa der von $\{a_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$ ($\{a_{\mu\nu}\}$ H. quadr.). Sei $f = \sum_1^\infty x_\mu \varphi_\mu$, wir setzen $f_N = \sum_1^N x_\mu \varphi_\mu$. Es ist $f_N \rightarrow f$, alle Sf_N sind sinnvoll, daher ist Sf sicher sinnvoll, wenn die Sf_N einen Limes haben (wegen der Abg.), d. h. wenn $|S(f_M - f_N)|^2 = |S(\sum_{N+1}^M x_\mu \varphi_\mu)|^2 \rightarrow 0$ für $N \rightarrow \infty$, $M \geq N$.

²⁵⁾ Die Möglichkeit, Operatoren zu addieren, ist insbesondere für die quantenmechanische Operatoralgebra von Bedeutung. Von einem einzigen H. O. kann man leichter Funktionen bilden; bei hypermax. können sogar, mit Hilfe der „Spektralform“ (d. h. der Z. d. E., A. E. Satz 36), alle Funktionen definiert werden, während $\bar{R}$ (A. E. Satz 37) gewisse Schwierigkeiten macht. Wir gehen hier darauf nicht näher ein.

Es ist

$$\left| S \left(\sum_{N+1}^M x_\mu \varphi_\mu \right) \right|^2 = \left| \sum_1^\infty \left(\sum_{N+1}^M a_{\mu\nu} x_\mu \right) \varphi_\nu \right|^2 = \sum_1^\infty \left| \sum_{N+1}^M a_{\mu\nu} x_\mu \right|^2 = \sum_1^\infty \left(\sum_{N+1}^M a_{\lambda\nu} \overline{a_{\lambda\mu}} x_\lambda \overline{x_\mu} \right),$$

und da alle Reihen $\sum_1^\infty a_{\lambda\nu} \overline{a_{\lambda\mu}} = A_{\lambda\mu}$ absolut konvergieren (A. E. Anm. 74),

$$= \sum_{N+1}^M A_{\lambda\mu} x_\lambda \overline{x_\mu} \leq \sum_{N+1}^M |A_{\lambda\mu}| \cdot |x_\lambda| \cdot |x_\mu|.$$

Nach einer in A. E., Beweis von Satz 46 sowie Anm. 60 angewandten Abschätzung (vgl. auch hier den Beweis von Satz 12) ist aber

$$\sum_1^\infty |A_{\lambda\mu}| \cdot |y_\lambda| |y_\mu| \leq \sum_1^\infty B_\mu |y_\mu|^2,$$

wenn $B_\mu = |A_{\mu\mu}| + 1 + \sum_0^{\mu-2} 2^{\sigma+1} |A_{\mu, \mu-\sigma-1}|$ (es ist $A_{\mu\lambda} = \overline{A_{\lambda\mu}}$) gesetzt wird. Aus

dem Spezialfalle $y_\mu = \begin{cases} 0 & \text{für } \mu \leq N \text{ und } \mu > M \\ x_\mu & \text{für } N < \mu \leq M \end{cases}$ dieser Formel folgt

$$\left| S \left(\sum_{N+1}^M x_\mu \varphi_\mu \right) \right|^2 \leq \sum_{N+1}^M |A_{\lambda\mu}| \cdot |x_\lambda| |x_\mu| \leq \sum_{N+1}^M B_\mu |x_\mu|^2.$$

Die gewünschte Konvergenz gegen 0 für $N \rightarrow \infty$, $M \geq N$, findet also gewiß statt,

wenn $\sum_1^\infty B_\mu |x_\mu|^2$ endlich ist.

Seien nun S, S' die max. H. O. von Satz 15, die nur für $f = 0$ beide Sinn haben. Sie seien etwa die abg. lin. H. O. von $\{a_{\mu\nu}\}$, $\{\varphi_1, \varphi_2, \dots\}$ und $\{b_{\mu\nu}\}$, $\{\psi_1, \psi_2, \dots\}$; wir bilden für beide die oben charakterisierten Folgen $B_1, B_2, \dots$ und $B'_1, B'_2, \dots$. Sei $c_1, c_2, \dots$ eine Folge reeller Zahlen mit $c_\mu^2 \geq B_\mu, B'_\mu$, und R, R' die abg. lin. H. O. von $\{c_\mu d_{\mu\nu}\}$, $\{\varphi_1, \varphi_2, \dots\}$ bzw. $\{c_\mu d_{\mu\nu}\}$, $\{\psi_1, \psi_2, \dots\}$, also Diagonaloperatoren.

Wenn für $f = \sum_1^\infty x_\mu \varphi_\mu = \sum_1^\infty x'_\mu \psi_\mu$ $Rf, R'f$ Sinn haben, so ist $\sum_1^\infty c_\mu^2 |x_\mu|^2, \sum_1^\infty c_\mu^2 |x'_\mu|^2$

endlich, also auch $\sum_1^\infty B_\mu |x_\mu|^2, \sum_1^\infty B'_\mu |x'_\mu|^2$ endlich — daher $Sf, S'f$ sinnvoll und $f = 0$.

Wenn U diejenige unitäre Transformation ist, für die $U\psi_\mu = \varphi_\mu$ ($\mu = 1, 2, \dots$) gilt, so ist offenbar $R' = URU^{-1}$. Damit ist alles bewiesen.

2. Der soeben bewiesene Satz kann auf alle H. O. übertragen werden. Wir beweisen dies in zwei Schritten.

Satz 17. Es gibt zu jedem nicht beschränkten Diagonaloperator R einen unitären U , so daß für kein $f \neq 0$ R und URU^{-1} beide Sinn haben.

Beweis: Sei R der abg. lin. H. O. von $\{a_\mu d_{\mu\nu}\}$, $\{\varphi_1, \varphi_2, \dots\}$; die a_μ dürfen nicht beschränkt sein²⁶⁾. Für ein gewisses Wertsystem a_μ , es heiße a'_μ , haben wir die Be-

²⁶⁾ Da sonst aus $f = \sum_1^\infty x_\mu \varphi_\mu$ (und alle $|a_\mu| \leq C$) folgen würde:

$$|Rf|^2 = \left| \sum_1^\infty a_\mu x_\mu \varphi_\mu \right|^2 = \sum_1^\infty |a_\mu|^2 \cdot |x_\mu|^2 \leq C^2 \cdot \sum_1^\infty |x_\mu|^2 = C^2 \cdot |f|^2, |Rf| \leq C \cdot |f|,$$

d. h. R wäre, entgegen der Annahme, beschränkt.

hauptung in Satz 16 bewiesen; allerdings nur für ein gewisses vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$, aber dieses spielt offenbar keine Rolle, da alle solchen Systeme durch Isomorphismen von $\mathfrak{H}$ ineinander überführbar sind. Wenn ferner statt $a_\mu = a'_\mu$ ($\mu = 1, 2, \dots$) nur $|a_\mu| \geq |a'_\mu|$ ($\mu = 2, 3, \dots$, für $\mu = 1$ nicht notwendig!) gilt, so trifft die Behauptung noch immer zu; denn der Definitionsbereich dieses R ist Teil von demjenigen des vorhergehenden, da ja aus der Endlichkeit von $\sum_1^\infty a_\mu^2 \cdot |x_\mu|^2$ die von $\sum_1^\infty a_\mu'^2 \cdot |x_\mu|^2$ folgt.

Nunmehr seien die a_μ ganz beliebig (aber unbeschränkt!). Wir zerlegen sie in eine Folge von Folgen $a_1^\varrho, a_2^\varrho, \dots$ ($\varrho = 1, 2, \dots$), derart daß stets $|a_\mu^\varrho| \geq |a'_\mu|$ ($\varrho = 1, 2, \dots, \nu = 2, 3, \dots$) ist²⁷⁾. Entsprechend teilen wir die $\varphi_1, \varphi_2, \dots$ in die Folgen $\varphi_1^\varrho, \varphi_2^\varrho, \dots$ ($\varrho = 1, 2, \dots$) ein. Die $\varphi_1^\varrho, \varphi_2^\varrho, \dots$ mögen die abg. lin. M. $\mathfrak{M}_\varrho$ aufspannen; die $\mathfrak{M}_\varrho$ sind offenbar paarweise orth. und spannen zusammen die abg. lin. M. $\mathfrak{H}$ auf. Man sieht, daß jedes $\mathfrak{M}_\varrho$ R reduziert; der darin liegende Teil von R sei R_ϱ ; offenbar ist R_ϱ (in $\mathfrak{M}_\varrho$ betrachtet) der abg. lin. H. O. von $\{a_\mu^\varrho \delta_{\mu\nu}\}, \{\varphi_1^\varrho, \varphi_2^\varrho, \dots\}$.

Also gibt es einen unitären Operator U_ϱ in $\mathfrak{M}_\varrho$, so daß nur für $f = 0$ R_ϱ und $U_\varrho R_\varrho U_\varrho^{-1}$ beide Sinn haben. Man konstruiert mühelos einen unitären Operator U in $\mathfrak{H}$, der durch jedes $\mathfrak{M}_\varrho$ reduziert wird, und dessen in $\mathfrak{M}_\varrho$ liegender Teil U_ϱ ist. Dann wird auch URU^{-1} durch $\mathfrak{M}_\varrho$ reduziert, und der darin liegende Teil von URU^{-1} ist $U_\varrho R_\varrho U_\varrho^{-1}$. Wenn nun für f R und URU^{-1} sinnvoll sind, so sind sie es auch für $P_{\mathfrak{M}_\varrho} f$, also auch $R_\varrho, U_\varrho R_\varrho U_\varrho^{-1}$ für $P_{\mathfrak{M}_\varrho} f$ sinnvoll. Also ist $P_{\mathfrak{M}_\varrho} f = 0, f$ zu $\mathfrak{M}_\varrho$ orth., also auch zu der von allen $\mathfrak{M}_\varrho$ aufgespannten abg. lin. M. $\mathfrak{H}$ — d. h. es ist $f = 0$.

Satz 18. Es gibt zu jedem nicht beschränkten H. O. R einen unitären Operator U , so daß für kein $f \neq 0$ R, URU^{-1} beide sinnvoll sind.

Beweis: Wegen Satz 17 genügt es zu zeigen: Es gibt einen unbeschränkten Diagonaloperator R' , der überall Sinn hat, wo R Sinn hat (aber eventuell andere Werte, also keine Forts.!). Da R max. Forts. besitzt, können wir uns auf max. H. O. R beschränken.

R' ist der abg. lin. H. O. eines $\{a_\mu \delta_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}$ (wobei die Unbeschränktheit der a_μ verlangt wird!), und es wird gewünscht: wenn Rf Sinn hat, $f = \sum_1^\infty x_\mu \varphi_\mu$, so soll $\sum_1^\infty a_\mu^2 |x_\mu|^2$ endlich sein — d. h. $\sum_1^\infty a_\mu^2 |(f, \varphi_\mu)|^2$. Von R' brauchen wir gar nicht zu sprechen, wir brauchen nur eine unbeschränkte Folge $a_1, a_2, \dots$ und ein vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$ mit dieser Eigenschaft. $\varphi_1, \varphi_2, \dots$ braucht übrigens nicht einmal vollst. zu sein; denn wir können dieses norm. orth. System durch Hinzunahme weiterer $\psi_1, \psi_2, \dots$ vollst. machen, und die entsprechenden $b_1, b_2, \dots$ gleich 0 wählen (sie spielen dann keine Rolle).

Es genügt hierzu ein norm. orth. System $\varphi_1, \varphi_2, \dots$ zu finden, so daß für alle f mit sinnvollem Rf

²⁷⁾ Etwa so: Sei $a_1^1 = a_1$, die $a_2^1, a_3^1, \dots$ eine $a'_2, a'_3, \dots$ absolut majorisierende Teilfolge der übrigen a_μ , die noch eine unbeschränkte Menge von a_μ übrigläßt. Sei a_1^2 das erste noch freie a_μ , die $a_2^2, a_3^2, \dots$ eine $a'_2, a'_3, \dots$ absolut majorisierende Teilfolge der übrigen a_μ , die noch eine unbeschränkte Menge von a_μ übrigläßt. Sei a_1^3 das erste noch freie a_μ , usw., usw. — So werden sicher alle $a_1, a_2, \dots$ genau aufgebraucht, und es ist offenbar $|a_\mu^\varrho| \geq |a'_\mu|$ für $\varrho = 1, 2, \dots, \nu = 2, 3, \dots$

$$|Rf|^2 \geq 8^\mu \cdot |(f, \varphi_\mu)|^2$$

ist, da hieraus

$$\sum_1^\infty 4^\mu \cdot |(f, \varphi_\mu)|^2 \leq \sum_1^\infty \frac{1}{2^\mu} |Rf|^2 = |Rf|^2$$

folgt (wir können also $a_\mu = 2^\mu$ setzen). Wir wollen eine hierfür hinreichende Bedingung angeben, die den Vorzug hat, daß in ihr nur die φ_μ vorkommen (und das unbestimmte f nicht). Sie lautet so: es sei $\varphi_\mu = Rg_\mu$, so daß auch R^2g_μ Sinn hat, und $|Rg_\mu| \geq C_\mu |g_\mu|$ ist (über C_μ verfügen wir gleich).

Dann ist nämlich $|Rg_\mu - C_\mu g_\mu| \leq |Rg_\mu| + C_\mu |g_\mu| \leq 2|Rg_\mu| = 2|\varphi_\mu| = 2$, und folglich

$$\begin{aligned} |Rf| &\geq |(Rf, \varphi_\mu)| = |(f, R\varphi_\mu)| \geq C_\mu |(f, \varphi_\mu)| - |(f, R\varphi_\mu - C_\mu \varphi_\mu)| \\ |(f, R\varphi_\mu - C_\mu \varphi_\mu)| &= |(f, R(Rg_\mu - C_\mu g_\mu))| = |(Rf, Rg_\mu - C_\mu g_\mu)| \\ &\leq |Rf| |Rg_\mu - C_\mu g_\mu| \leq 2|Rf|, \end{aligned}$$

also $3|Rf| \leq C_\mu |(f, \varphi_\mu)|$, $|Rf|^2 \geq \frac{C_\mu^2}{3} |(f, \varphi_\mu)|^2$ — daher genügt $C_\mu = 3\sqrt{8^\mu}$. Ein solches norm. orth. System $\varphi_1, \varphi_2, \dots$ brauchen wir also.

Wenn der H. O. R nicht hypermax. ist, so gilt es eine abg. lin. M. die R reduziert, und für die der dort liegende Teil von R gleich $\bar{R}$ oder $-\bar{R}$ ist — wenn wir unsere Aufgabe hierfür lösen, ist sie auch für R gelöst. Wir brauchen also nur hypermax. R zu betrachten, und $\bar{R}$ ($-\bar{R}$ ist dann miterledigt). Für diese Operatoren müssen wir die Existenz einer Folge $g_1, g_2, \dots$ mit den folgenden Eigenschaften nachweisen: Alle Rg_μ , R^2g_μ sind sinnvoll, die Rg_μ zueinander orth., $|Rg_\mu| \geq C_\mu \cdot |g_\mu|$, und $g_\mu \neq 0$ ($|\varphi_\mu| = |Rg_\mu| = 1$ brauchen wir nicht zu verlangen, wir können die g_μ nachher so normieren). Dabei ist $C_\mu = 3\sqrt{8^\mu}$.

Für $\bar{R}$ soll dies im Anhang erfolgen, für hypermax. R machen wir es jetzt. Sei $E(\mu)$ die Z. d. E. die zu R gehört (A. E. Satz 36); da R nicht beschränkt ist, gibt es eine Folge $\lambda_1, \lambda_2, \dots$ mit $\lambda_1 < \lambda_2 < \dots$ oder $\lambda_1 > \lambda_2 > \dots$, $\lambda_\mu \rightarrow \infty$ bzw. $\rightarrow -\infty$, und $E(\lambda_1) \neq E(\lambda_2)$, $E(\lambda_2) \neq E(\lambda_3), \dots$ (vgl. § 5, Kap. III). Wie dort, brauchen wir nur den ersten Fall zu betrachten, da der zweite ebenso zu behandeln ist. Die λ_μ können wir gleich mit $\lambda_\mu \geq C_\mu$ wählen. Wieder sei $E(\lambda_{\mu+1}) - E(\lambda_\mu) = P_{\mathfrak{M}_\mu}$, $\mathfrak{M}_\mu \neq (0)$, und die $\mathfrak{M}_\mu$ sind paarweise orth. Wenn g zu $\mathfrak{M}_\mu$ gehört, so ist

$$\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)g|^2 = \int_{\lambda_\mu}^{\lambda_{\mu+1}} \lambda^2 d|E(\lambda)g|^2 \leq \lambda_{\mu+1}^2 \int_{\lambda_\mu}^{\lambda_{\mu+1}} d|E(\lambda)g|^2 \leq \lambda_{\mu+1}^2 \cdot |g|^2,$$

also endlich — d. h. Rg sinnvoll. Da $\mathfrak{M}_\mu$ R reduziert (vgl. § 5, Kap. III) gehört Rg wieder zu $\mathfrak{M}_\mu$, und auch R^2g ist daher sinnvoll. Schließlich ist

$$\begin{aligned} (g, Rg) &= \int_{-\infty}^{\infty} \lambda d|E(\lambda)g|^2 = \int_{\lambda_\mu}^{\lambda_{\mu+1}} \lambda d|E(\lambda)g|^2 \geq C_\mu \int_{\lambda_\mu}^{\lambda_{\mu+1}} d|E(\lambda)g|^2 \\ &= C_\mu \int_{-\infty}^{\infty} d|E(\lambda)g|^2 = C_\mu \cdot |g|^2, \end{aligned}$$

$$C_\mu \cdot |g|^2 \leq (g, Rg) \leq |g| \cdot |Rg|, \quad |Rg| \geq C_\mu \cdot |g|.$$

Daher können wir für g_μ irgendein Element $\neq 0$ von $\mathfrak{M}_\mu$ wählen. Da die Rg_μ in den $\mathfrak{M}_\mu$ liegen, sind sie paarweise orth., und daß alles andere erfüllt ist, sahen wir schon.

Man beachte, daß sich wieder die am Ende der Einleitung und des Kap. V angezeigte Erscheinung zeigt: Satz 18 gilt für alle unbeschränkten H. O. R ohne Ausnahme, nur der unitäre Operator U muß zweckmäßig gewählt werden — wenn man also hierin eine Pathologie sieht, so hat man dieselbe bei den unitären Operatoren zu suchen.

Anhang: Eigenschaften von $\bar{R}$.

1. Wir müssen zwei Eigenschaften des $\bar{R}$ aus A. E. Satz 37 nachweisen, die im § 5, Kap. III und beim Beweise von Satz 18 verwendet wurden. Hauptsächlich um der ersteren willen werden wir zunächst eine weitere Realisation von $\bar{R}$ angeben, die schon in A. E. Anhang IV erwähnt wurde.

Sei Ω das Intervall $0 \leq x < \infty$, wir bilden seinen Funktionenraum (vgl. A. E. Einleitung I, III und Anhang I): die Menge aller (nach Lebesgue meßbaren) Funktionen $f(x)$ des Intervalles $0 \leq x < \infty$ mit endlichem $\int_0^\infty |f(x)|^2 dx$ — nach A. E. Anhang I ist dies ein Hilbertscher Raum, isomorph $\mathfrak{H}$. Er heiße $\mathfrak{H}^{**}$.

$\mathfrak{A}$ sei die Menge aller stetigen und bis auf endlich viele Knicke überall stetig differenzierbaren Funktionen $f(x)$ von $\mathfrak{H}^{**}$, mit endlichem $\int_0^\infty |f'(x)|^2 dx$ und $f(0) = 0$. Der für diese $f(x)$ definierte Operator $i \frac{d}{dx} \cdots$ heiße S . S ist offenbar lin., ferner ein H. O., denn es ist

$$(f, Sg) - (Sg, f) = -i \int_0^\infty \{f(x) \bar{g}'(x) + f'(x) \bar{g}(x)\} dx = -i [f(x) \bar{g}(x)]_0^\infty = 0^{28}),$$

und $\mathfrak{A}$ ist überall dicht ²⁹⁾. Wir werden zeigen, daß $\tilde{S}$ (in einem geeigneten vollst. norm. orth. System $\varphi_1, \varphi_2, \dots$, vgl. A. E. Satz 37) gleich $\bar{R}$ ist. Hierzu müssen wir seine Cayleysche Transformierte U bestimmen.

Betrachten wir das durch $S \pm i \cdot 1$ vermittelte Bild von $\mathfrak{A}$. Es ist

$$(S \pm i \cdot 1)f(x) = i(f'(x) \pm f(x)),$$

also wenn dies $= \varphi(x)$ ist,

$$i(f'(x) \pm f(x)) = \varphi(x), \quad (e^{\pm x} f(x))' = -i e^{\pm x} \varphi(x), \quad f(x) = -i e^{\mp x} \int_0^x e^{\pm y} \varphi(y) dy$$

(f , weil $f(0) = 0$ sein muß). Für das Vorzeichen — sehen wir insbesondere: $\int_0^\infty e^{-y} \varphi(y) dy$

²⁸⁾ Hierzu müssen wir noch $\lim_{x \rightarrow \infty} f(x) \bar{g}(x) = 0$ beweisen. Wegen der feststehenden Konvergenz des Integrals ist $\lim_{x \rightarrow \infty} f(x) \bar{g}(x)$ vorhanden; wäre er $= a \neq 0$, so wäre $\lim_{x \rightarrow \infty} (|f(x)|^2 + |g(x)|^2) \geq |a| > 0$, was mit der

Endlichkeit von $\int_0^\infty |f(x)|^2 dx, \int_0^\infty |g(x)|^2 dx$ unvereinbar ist.

²⁹⁾ Da A eine lin. M. ist, genügt es zu zeigen, daß es eine Funktionenmenge zu Häufungspunkten hat, die die Abg. lin. M. $\mathfrak{H}^{**}$ aufspannt. Eine solche ist nach A. E. Anhang I (Diskussion der Bedingung C) die Menge aller Funktionen $f(x)$, die in einer aus endlich vielen offenen Intervallen bestehenden Menge $= 1$ und sonst $= 0$ sind. (Man wähle nämlich für das dort genannte Umgebungssystem $I_1, I_2, \dots$ die Menge aller rationalendigen Intervalle.) Daß aber solche Funktionen durch überall stetig differenzierbare $f(x)$, mit $f(0) = 0$ und endlichem $\int_0^\infty |f(x)|^2 dx$ und $\int_0^\infty |f'(x)|^2 dx$, im Mittel beliebig gut approximierbar sind, ist klar.

konvergiert, und wenn es $\neq 0$ ist, so ist $\lim_{x \rightarrow \infty} f(x) = \infty$; da $\int_0^{\infty} |f(x)|^2 dx$ endlich ist, darf dies nicht der Fall sein, also ist $(S - i \cdot 1)f(x)$ immer zu e^{-x} (das ja zu $\mathfrak{H}^{**}$ gehört!) orth. Wenn $(S + i \cdot 1)f(x) = \varphi(x)$ ist, so ist

$$U\varphi(x) = (S - i \cdot 1)f(x) = \varphi(x) - 2if(x) = \varphi(x) - 2e^{-x} \int_0^x e^y \varphi(y) dy.$$

Wenn nun $\varphi(x)$ die Form $e^{-x}p(x)$, $p(x)$ ein Polynom, hat, so gehört es sicher zu $\mathfrak{H}^{**}$; das aus $(R + i \cdot 1)f(x) = \varphi(x)$ bestimmte

$$f(x) = -ie^{-x} \int_0^x p(y) dy = e^{-x}q(x), \quad q(x) = -i \int_0^x p(y) dy$$

hat also dieselbe Form, gehört auch zu $\mathfrak{H}^{**}$, also sogar zu $\mathfrak{A}$ — also ist $U\varphi(x)$ sinnvoll:

$$U\varphi(x) = e^{-x}p(x) - 2e^{-x} \int_0^x p(y) dy = e^{-x}r(x), \quad r(x) = p(x) - 2 \int_0^x p(y) dy,$$

und hat dieselbe Form. Wir bilden daher sukzessive:

$$\varphi_0(x) = e^{-x}, \quad \varphi_1(x) = U\varphi_0(x), \quad \varphi_2(x) = U\varphi_1(x), \dots;$$

hier ist $\varphi_n(x) = e^{-x}p_n(x)$,

$$p_0(x) = 1, \quad p_{n+1}(x) = p_n(x) - 2 \int_0^x p_n(y) dy.$$

Somit ist $p_n(x)$ ein Polynom von genau n -tem Grade.

Alle $\varphi_k(x)$ ($k = 1, 2, \dots$) haben die Form $(S - i \cdot 1)f(x)$, sind also zu $\varphi_0(x) = e^{-x}$ orth., also wenn wir das längentreue U^m ($m = 0, 1, 2, \dots$) anwenden: $\varphi_{m+k}(x)$ zu $\varphi_m(x)$. D. h. $\varphi_m(x), \varphi_n(x)$ sind orth. für $m < n$, also für alle $m \neq n$. Weiter ist, wegen der Längentreue von U ,

$$|\varphi_m(x)|^2 = |\varphi_0(x)|^2 = \int_0^{\infty} e^{-2x} dx = \frac{1}{2}.$$

Die $\sqrt{2}\varphi_m(x)$ ($n = 0, 1, 2, \dots$) bilden somit ein norm. orth. System.

Die $\sqrt{2}p_m(x)$ sind also ein System von Polynomen der Grade $m = 0, 1, 2, \dots$, die mit der „Belegung“ e^{-2x} im Intervalle $0, \infty$ norm. orth. sind, also die $p_n(\frac{1}{2}x)$ mit der „Belegung“ e^{-x} — folglich müssen es die Laguerreschen Polynome sein. Diese aber bilden bekanntlich ein vollst. System³⁰⁾, somit sind die $\sqrt{2}\varphi_0(x), \sqrt{2}\varphi_1(x), \sqrt{2}\varphi_2(x), \dots$ auch ein vollst. norm. orth. System.

Der längentreue Operator U ist im vollst. norm. orth. System $\sqrt{2}\varphi_n(x)$ ($n = 1, 2, 3, \dots$) stets sinnvoll, also ist sein Definitionsbereich (eine abg. lin. M.) ganz $\mathfrak{H}^{**}$. Es ist $U\varphi_n(x) = \varphi_{n+1}(x)$, also ist U der Operator $\bar{U}$ (A. E. Satz 37), und folglich seine Cayleysche Transformierte $\bar{S}$ gleich $\bar{R}$. — Damit haben wir gezeigt:

Satz 19. $\mathfrak{H}^{**}$ sei der soeben beschriebene Hilbertsche Funktionenraum der $f(x)$ von $0 \leq x < \infty$, $\mathfrak{A}$ die Menge aller stetigen und bis auf endlich viele Knicke überall stetig differentiierbaren $f(x)$ von $\mathfrak{H}^{**}$ mit $f(0) = 0$. S sei der in $\mathfrak{A}$ definierte Operator $i \frac{d}{dx} \dots$; dann ist der abg. lin.

³⁰⁾ Für den Beweis dieser Tatsache vgl. z. B. das Lehrbuch: *Courant-Hilbert, Methoden der mathematischen Physik*, Berlin 1926.

H. O. $\tilde{S}$ im soeben beschriebenen vollst. norm. orth. System dem Operator $\bar{R}$ (A. E. Satz 37) gleich.

2. Wir beweisen nun die gewünschten Eigenschaften von $\bar{R}$. $\mathfrak{A}, S$ seien, wie bisher. $\mathfrak{A}'$ sei die Menge aller Funktionen aus $\mathfrak{A}$, für die $f(0) = f(1) = f(2) = \dots = 0$ ist, $\mathfrak{A}'$ ist eine lin. M. und überall dicht³¹⁾. Das auf $\mathfrak{A}'$ eingeschränkte S heie S' , nach dem soeben Gesagten ist es ein lin. H. O., S ist Forts. von S' , also $\bar{R} = \tilde{S}$ von $\tilde{S}'$.

Nun sei $\mathfrak{M}_n$ die Menge aller Funktionen $f(x)$ von $\S^{**}$, für die $f(x) = 0$ für $x \leq n$ und $x \geq n + 1$ gilt. Offenbar ist jedes $\mathfrak{M}_n$ eine abg. lin. M., je zwei $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ sind orth., und alle zusammen spannen die abg. lin. M. $\S^{**}$ auf. $P_{\mathfrak{M}_n}f(x)$ ist $\begin{cases} f(x) & \text{für } n < x < n + 1 \\ 0 & \text{sonst} \end{cases}$, also gehört mit $f(x)$ auch $P_{\mathfrak{M}_n}f(x)$ zu $\mathfrak{A}'$ (für $\mathfrak{A}$ gilt dies nicht!). Hieraus folgert man leicht, daß die Bedingungen des Reduziertwerdens durch $\mathfrak{M}_n$ von S' erfüllt werden (A. E. Def. 13), also auch von $\tilde{S}'$; also wird $\tilde{S}'$ von allen $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ reduziert. Die im § 5, Kap. III benützte Eigenschaft von $\bar{R}$ ist damit bewiesen (der gesuchte Operator ist S')³²⁾.

Die andere Aufgabe war: ein Funktionensystem $g_1, g_2, \dots$ (alle $\neq 0$) zu finden, so daß die $\bar{R}g_\mu, \bar{R}^2g_\mu$ sinnvoll sind, die $\bar{R}g_1, \bar{R}g_2, \dots$ zueinander orth. sind, und $|\bar{R}g_\mu| \geq C_\mu \cdot |g_\mu|$ ($C_\mu = 3\sqrt{8^\mu}$). Dies leisten, wie man sich leicht überlegt, z. B. die Funktionen (aus $\mathfrak{A}$!)

$$g_\mu(x) = \begin{cases} [4(x - \mu)(\mu + 1 - x)]^{D_\mu} & \text{für } \mu < x < \mu + 1 \\ 0 & \text{sonst} \end{cases} \quad (D_\mu \geq 1),$$

wenn die D_μ groß genug gewählt werden.

³¹⁾ Dies ergibt sich durch dieselben Überlegungen, wie sie in Anm. ²⁹⁾ für $\mathfrak{A}$ angestellt wurden.
³²⁾ Das nach A. E. Satz 38 irreduzible $\bar{R}$ ist also einem reduzierbaren k. u. äquivalent! — Daß für jedes $\mathfrak{M}_n$ der darin liegende Teil von $\tilde{S}'$ nicht beschränkt ist, folgt z. B. aus dem im Text folgenden Beispiel; denn g_n gehört zu $\mathfrak{M}_n$ und zu $\mathfrak{A}'$, und wenn D_n groß genug gewählt wird, so ist $(g_n, \bar{R}g_n): |g|^2$ auch beliebig groß.

Über einen Hilfssatz der Variationsrechnung.

1. In der vorhergehenden Arbeit von Herrn A. HAAR kommt der folgende, vom genannten Verfasser formulierte und auf seine Anregung zuerst von Herrn T. RADÓ bewiesene Satz zur Anwendung, der auch an und für sich nicht ohne Interesse ist¹⁾:

x, y, z sei ein räumliches (Kartesisches) Koordinatensystem. $\mathfrak{R}$ sei ein konvexer (und beschränkt abgeschlossener) Bereich in der x, y -Ebene, F eine durch die Gleichung $z = f(x, y)$ definierte räumliche Fläche über $\mathfrak{R}$ (d. h. x, y soll in $\mathfrak{R}$ liegen, die Funktion $f(x, y)$ sei stetig). Δ sei eine feste positive Zahl.

Wenn für eine Ebene $E, z = ax + by + c$, ein Teilbereich $\mathfrak{A}$ von $\mathfrak{R}$ existiert, so daß für alle x, y , die am Rande von $\mathfrak{A}$ liegen, die darüberliegenden Punkte von E und F zusammenfallen (d. h. $f(x, y) = ax + by + c$ ist), so soll auch im Inneren von $\mathfrak{A}$ dieses Zusammenfallen von E und F stattfinden²⁾. Keine Ebene $E, z = ax + by + c$, deren „maximale Steigung“ $\sqrt{a^2 + b^2} > \Delta$ ist, soll den Rand von F ³⁾ in drei oder mehr⁴⁾ Punkten treffen.

Dann erfüllt $f(x, y)$ in ganz $\mathfrak{R}$ die Lipschitzsche Bedingung:

$$|f(x, y) - f(x', y')| \leq \Delta \sqrt{(x - x')^2 + (y - y')^2}.$$

Im folgenden soll ein Beweis dieses Satzes gegeben werden, welcher, wie wir glauben, kürzer und einfacher ist, als derjenige von RADÓ.

¹⁾ Vgl. A. HAAR: „Über das Plateausche Problem“, Math. Annalen. Bd. 97, S. 124, ferner den vorhergehenden Artikel in dieser Zeitschrift, S. 1–27, sowie T. RADÓ: „Geometrische Betrachtungen über reguläre Variationsprobleme“, Acta Litt. ac Scient., Univ. Szeged, Bd. 2, S. 228–253. Die Entstehung, Bedeutung und Literatur dieses Satzes wird an diesen Stellen genauer erörtert.

²⁾ Diese Prämisse darf dahin abgeschwächt werden, daß im Inneren von $\mathfrak{A}$ nicht stets $f(x, y) > ax + by + c$ sein darf und auch nicht stets $f(x, y) < ax + by + c$. Wäre nämlich einmal $f(x, y) = ax + by + c$, also $\geq ax + by + c$, so sei $\mathfrak{A}'$ die Menge der x, y in $\mathfrak{A}$ mit $f(x, y) >$ bzw. $< ax + by + c$ mit ihrem Rand. Der Rand von $\mathfrak{A}'$ ist, soweit er nicht am Rande von $\mathfrak{A}$ liegt, die Menge der gemeinsamen Häufungspunkte der x, y in $\mathfrak{A}$ mit $f(x, y) >$ bzw. $< ax + by + c$ und der x, y in $\mathfrak{A}$ mit $f(x, y) \leq$ bzw. $\geq ax + by + c$. Also ist dort $f(x, y) = ax + by + c$, am Rande von $\mathfrak{A}$ gilt dies nach Annahme auch. Verletzt also $\mathfrak{A}$ die alte Prämisse, so verletzt $\mathfrak{A}'$ auch die neue. — Übrigens wird beim Beweise im wesentlichen nur diese, weniger weitgehende, Prämisse ausgenutzt.

³⁾ Den Rand von F bilden diejenigen Punkte x, y, z ($z = f(x, y)$), die über den Randpunkten x, y von $\mathfrak{R}$ liegen.

⁴⁾ Es darf, ohne wesentliche Beeinträchtigung des auszuführenden Beweises, hinzugefügt werden: nicht auf einer Gerade liegenden Punkten.

Von nun an also denken wir uns $\mathfrak{R}$, F , $f(x, y)$, Δ fest und die Prämissen des Satzes erfüllend vorgegeben, es gilt die Lipschitzsche Bedingung als gültig nachzuweisen.

2. Betrachten wir irgendeine Ebene E , $z = ax + by + c$, die sowohl Punkte über, als auch solche unter F hat (d. h. $ax + by + c \leq f(x, y)$ komme beides vor). Die in die x, y -Ebene projizierte Menge der ersteren heiße $\mathfrak{A}$, die der letzteren $\mathfrak{B}$. Jede Komponente von $\mathfrak{A}$ und von $\mathfrak{B}$ hat (wenn ihr Rand hinzugefügt wird) lauter innere Punkte mit $f(x, y) \neq ax + by + c$, also muß (nach der Prämisse des Satzes) dies auch am Rande vorkommen. Nun ist jeder Randpunkt auch einer von $\mathfrak{A}$ bzw. $\mathfrak{B}$, also Häufungspunkt der x, y mit $f(x, y) <$ bzw. $> ax + by + c$, also gilt dort $f(x, y) \leq$ bzw. $\geq ax + by + c$. Liegt er nicht am Rande von $\mathfrak{R}$, so ist er auch Häufungspunkt von $\mathfrak{R} - \mathfrak{A}$ bzw. $\mathfrak{R} - \mathfrak{B}$, also der x, y mit $f(x, y) \geq$ bzw. $\leq ax + by + c$, so daß dann dort $f(x, y) = ax + by + c$ sein wird. Die obengenannten Randpunkte liegen daher am Rande von $\mathfrak{R}$, und in ihnen ist $f(x, y) <$ bzw. $> ax + by + c$, d. h. sie liegen auch in $\mathfrak{A}$ bzw. $\mathfrak{B}$. Der Durchschnitt von $\mathfrak{A}$ bzw. $\mathfrak{B}$ mit dem Rande von $\mathfrak{R}$ heiße $\mathfrak{A}_1$ bzw. $\mathfrak{B}_1$, dann haben wir gezeigt, jede Komponente von $\mathfrak{A}$ bzw. $\mathfrak{B}$ enthält Punkte aus $\mathfrak{A}_1$ bzw. $\mathfrak{B}_1$.

Sei nun die maximale Steigung von E $\sqrt{a^2 + b^2} > \Delta$. Dann darf auf dem Rande von $\mathfrak{R}$ nicht dreimal $f(x, y) = ax + by + c$ werden, d. h. höchstens zwei Punkte dieses Randes gehören weder zu $\mathfrak{A}_1$ noch zu $\mathfrak{B}_1$. Dabei sind $\mathfrak{A}_1$, $\mathfrak{B}_1$ offenbar (als Teilmengen dieses Randes) relativ offen. Da $\mathfrak{A}$, $\mathfrak{B}$ nicht leer sind, also Komponenten besitzen, ist weder $\mathfrak{A}_1$ noch $\mathfrak{B}_1$ leer. Schließlich ist der Rand von $\mathfrak{R}$ offenbar eine konvexe (geschlossene) Kurve. Hieraus folgert man sofort, daß $\mathfrak{A}_1$, $\mathfrak{B}_1$ zwei einfache Bögen dieser Kurve sind, die in zwei Punkten zusammenstoßen (in diesen letzteren, und nur in ihnen, gilt dann $f(x, y) = ax + by + c$). Also sind $\mathfrak{A}_1$, $\mathfrak{B}_1$ beide zusammenhängend. Da jede Komponente von $\mathfrak{A}$ bzw. $\mathfrak{B}$ Punkte aus $\mathfrak{A}_1$ bzw. $\mathfrak{B}_1$ enthält, muß sie $\mathfrak{A}_1$ bzw. $\mathfrak{B}_1$ sogar ganz umfassen. Zwei verschiedene Komponenten von $\mathfrak{A}$ bzw. $\mathfrak{B}$ hätten darum gemeinsame Punkte, also kann es keine solche geben, d. h. $\mathfrak{A}$ ist zusammenhängend und $\mathfrak{B}$ auch.

3. Sei g eine Gerade mit einer Steigung $> \Delta$. Jede Ebene E , die g enthält, hat somit eine maximale Steigung, die erst recht $> \Delta$ ist, also gilt das in 2. bewiesene für jedes solche E . Wir können das Koordinatensystem derart um die z -Achse herum drehen, daß die Projektion von g in die x, y -Ebene die Richtung der x -Achse aufweist, und zwar die positive Richtung diejenige ist, in der g hinaufsteigt. Ferner genügt eine Parallelverschiebung der x, y -Ebene (das ist eine Koordinatentransformation $z \rightarrow z + \text{Konstans}$), um zu erreichen, daß g durch den Koordinaten-Nullpunkt geht. Daher wollen wir uns das Koordinaten-

system von vornherein so gewählt denken. Seien nun P, Q, R drei Punkte auf g , in steigender Reihenfolge, d. h. ihre Projektionen p, q, r auf der x -Achse mögen ein in dieser Reihenfolge wachsendes x aufweisen. Wir wollen zeigen: wenn g in P, R über den entsprechenden Punkten von F liegt, so liegt Q nicht darunter, und wenn es in P, R unter F liegt, so liegt P nicht darüber. Da das zweite durch Vorzeichenwechsel von z auf erstere zurückführbar ist, betrachten wir nur den ersten Fall. Liege also in P, R g über F , in Q aber darunter: es gilt dies ad absurdum zu führen.

Sei θ ein Winkel $> -\frac{\pi}{2}, < \frac{\pi}{2}$. Wir ziehen in der y, z -Ebene eine Gerade durch den Nullpunkt, die mit der positiven y -Achse den Winkel θ einschließt, und durch diese Gerade und durch g (beide schneiden sich im Nullpunkte!) legen wir eine Ebene, E_θ . E_θ hat stets die Form $z = ax + by + c$, für $\theta \rightarrow \pm \frac{\pi}{2}$ strebt es aber gegen die x, z -Ebene, die diese Form nicht zuläßt. Die maximale Steigung von E_θ ist, wie wir schon bemerkten, stets $> \Delta$, wenn wir also für jedes E_θ die obigen Mengen $\mathfrak{A}_\theta, \mathfrak{B}_\theta$ bilden, so sind beide zusammenhängend. Für $\theta \rightarrow \frac{\pi}{2}$ wächst rechts von der x -Achse (in der x, y -Ebene, d. h. für $y > 0$) jede z -Koordinate auf E_θ gegen $+\infty$, d. h. jeder Punkt der x, y -Ebene mit $y > 0$ liegt schließlich in $\mathfrak{B}_\theta$; analog liegt jeder mit $y < 0$ schließlich in $\mathfrak{A}_\theta$. Für $\theta \rightarrow -\frac{\pi}{2}$ ist es umgekehrt. Über der x -Achse selbst ($y = 0$) enthält jedes E_θ g , hier sind also die Verhältnisse von θ unabhängig: p, r gehören zu $\mathfrak{B}_\theta$, q aber zu $\mathfrak{A}_\theta$.

Für jedes θ gibt es also einen Polygonzug, der p, r in $\mathfrak{B}_\theta$ verbindet, die Arcusvariation längs desselben, um den Mittelpunkt q genommen, ist offenbar $\equiv \pi \dots \text{mod } (2\pi)$, (da p, q, r in dieser Reihenfolge auf einer Geraden liegen). Wäre sie für zwei verschiedene solche Polygonzüge verschieden, so bildeten diese zusammen ein Polygon, dessen Arcusvariation um q herum $\neq 0$ ist, d. h. in dessen Inneren r liegt. Also: $\mathfrak{A}_\theta$ hat Punkte im Inneren dieses Polygons, und keine darauf (es liegt ja ganz in $\mathfrak{B}_\theta$), da es zusammenhängend ist, liegt es somit ganz im Inneren desselben. Also erst recht im Inneren von $\mathfrak{R}$, entgegen der bewiesenen Tatsache, daß es Punkte am Rande von $\mathfrak{R}$ haben muß. Somit ist die oben genannte Arcusvariation eine eindeutig festgelegte Zahl, etwa φ_θ . Es ist $\varphi_\theta \equiv \pi \dots \text{mod } (2\pi)$.

Wenn ein Polygonzug von p nach r in $\mathfrak{A}_\theta$ liegt, so liegt er auch für alle genügend nahe bei θ liegenden θ' in $\mathfrak{A}_{\theta'}$, so daß $\varphi_{\theta'} = \varphi_\theta$ sein muß. Somit ist φ_θ überhaupt konstant.

Da p, r zu $\mathfrak{B}_0$ gehören, gehören auch genügend kleine Umgebungen von p, r dazu, insbesondere zwei geeignete, von p, r in die Halbebene

$y > 0$ hineinführende Strecken π, ϱ . Für $\theta > 0$, $\theta \rightarrow \frac{\pi}{2}$ steigt E_θ in der Halbebene $y > 0$, so daß π, ϱ für alle $\theta > 0$ erst recht in $\mathfrak{B}_\theta$ liegt. Die Verbindungslinie σ der freien Enden von π, ϱ liegt ganz in der Halbebene $y > 0$, für $\theta \rightarrow \frac{\pi}{2}$ liegt sie also schließlich auch in $\mathfrak{B}_\theta$. π, σ, ϱ ist also schließlich ein in $\mathfrak{B}_\theta$ liegender Polygonzug von p nach r , da er ganz in der genannten Halbebene liegt, ist seine Arcusvariation um q herum π . D. h.: für $\theta \rightarrow \frac{\pi}{2}$ wird schließlich $\varphi_\theta = \pi$. Ebenso beweist

man, daß für $\theta \rightarrow -\frac{\pi}{2}$ schließlich $\varphi_\theta = -\pi$ wird.

Dies widerspricht der Konstanz von φ_θ , womit das vorhin Behauptete bewiesen ist.

4. Seien nun P, Q wieder zwei Punkte auf g , so daß g von P nach Q wächst, und es liege P über und Q unter F . Die Projektionen von P, Q sind p, q auf der x -Achse, die somit in $\mathfrak{R}$ eindringt, daher die Randkurve von $\mathfrak{R}$ zweimal trifft: in r, s ($\vec{r}, \vec{s}$ sei die positive Richtung). Über r, s sollen bzw. die Punkte R, S von g liegen, und die bzw. Punkte R', S' von F (dies sind Randpunkte von F !). Nach dem vorhin Bewiesenen liegt R nicht unter und S nicht über F , d. h. R ist gleich oder über R' , und S gleich oder unter S' . Ersetzen wir nun die Gerade R, S , d. h. g , durch die Gerade R', S' , sie heiße g' , so nimmt demnach die Steigung nicht ab, sie bleibt also $> \Delta$. Sei T irgendein dritter Randpunkt von F , die eine durch R', S', T gelegte Ebene enthält dann g' , hat also erst recht maximale Steigung $> \Delta$. Trotzdem enthält sie drei Randpunkte von F (R', S', T), entgegen der Prämisse des Satzes.

Daher ist die eingangs beschriebene Lage von P, Q auf g (relativ zu F) unmöglich. Es ist aber auch unmöglich, daß P, Q beide auf F liegen. Denn ein genügend kleines Heben von g auf dem Lot aus P auf die x, y -Ebene, verbunden mit einem genügend kleinen Senken auf dem Lot aus Q auf die x, y -Ebene, würde die Steigung $> \Delta$ belassen, aber doch bewirken, daß der über P bzw. Q liegende Punkt des geänderten g über P bzw. unter Q liegt, d. h. über bzw. unter F .

Somit kann eine Gerade g , deren Steigung $> \Delta$ ist, F nicht in zwei Punkten treffen. Sind $x, y, f(x, y)$ und $x', y', f(x', y')$ irgend zwei verschiedene Punkte von F , so ist die Steigung ihrer Verbindungsgeraden $\frac{|f(x, y) - f(x', y')|}{\sqrt{(x - x')^2 + (y - y')^2}}$. Diese muß also $\leq \Delta$ sein, d. h. die behauptete Lipschitzsche Bedingung muß gelten (fallen x, y und x', y' zusammen, so ist sie trivial).

Damit ist das gewünschte Ziel erreicht.

ÜBER FUNKTIONEN VON FUNKTIONALOPERATOREN

Einleitung.

1. Den Gegenstand der vorliegenden Arbeit bilden gewisse Sätze über die funktionelle Abhängigkeit vertauschbarer (beschränkter) Funktionaloperatoren. Dieselben hängen mit Fragen zusammen, welche bereits in einer früheren Arbeit des Verfassers angeschnitten wurden¹⁾, deren vollständige Diskussion dort aber (aus räumlichen Rücksichten, da die genannte Arbeit auch anderen Gegenständen gewidmet war) nicht erfolgte. Hier soll eine erschöpfende Untersuchung und Erledigung erfolgen, im Verein mit einer (dadurch nahegelegten) vollständigen Entwicklung des Funktionsbegriffes für Funktionaloperatoren. Es handelt sich um den Beweis des folgenden Satzes: Zwei (oder sogar beliebig viele) vertauschbare²⁾ Hermitesche Operatoren können immer als Funktionen eines einzigen dargestellt werden. In A, a. a. O. Anm. ¹⁾, wurde nämlich nur gezeigt, daß sie in einem gewissen Sinne Limites von Polynomen desselben sind, da der allgemeine Begriff einer Funktion eines Operators dort nicht entwickelt wurde. Der oben genannte Satz soll hier übrigens in einer noch etwas verschärften und präzisierten Form bewiesen werden, auf die wir noch zu sprechen kommen.

2. Unter den für die folgenden Betrachtungen wesentlichen Begriffsbildungen stehen an erster Stelle diejenigen des Hilbertschen Raumes und seiner linearen Operatoren. Es ist wohl nicht nötig, dieselben des Näheren zu entwickeln, aber da wir bei ihrer Untersuchung eine geometrische, besonders für unsere Zwecke zugeschnittene, Terminologie verwenden werden, sei gesagt, daß wir uns an die diesbezüglichen Ausführungen und Begriffsbildungen von E und A anschließen³⁾. Wie dort, nennen wir den (abstrakten) Hilbertschen Raum $\mathfrak{H}$, den Ring seiner beschränkten Operatoren $\mathbf{B}$. Ferner wird von den Ringen in $\mathbf{B}$ immer wieder die Rede sein: das sind diejenigen Teilmengen von $\mathbf{B}$, welche mit allen Elementen A, B auch $\alpha A, A^{*4)}$, $A + B$, AB enthalten, und außerdem eine gewisse Abgeschlossenheits-Eigenschaft besitzen⁵⁾. Jede Teilmenge $\mathbf{M}$ von $\mathbf{B}$ ist in gewissen Ringen enthalten, u. a. in einem kleinsten: das ist der von ihr erzeugte Ring $\mathbf{R}(\mathbf{M})^{6)}$. Besteht insbesondere $\mathbf{M}$ aus einem einzigen Element A , so können wir aus A den Ring $\mathbf{R}(A)$ erzeugen. Dieser besteht

übrigens aus allen Polynomen $p(A)$ von A ($p(x)$ durchläuft die Polynome mit $p(0) = 0$), sowie den Limites aller sogenannten stark doppelkonvergenten Folgen derselben⁷⁾ — wenigstens, wenn A Hermitesch oder zumindest normal ist⁸⁾).

Eine Menge M (in B) ist Abelsch, wenn alle ihre Elemente miteinander und mit ihren $*$ vertauschbar sind; bei Mengen, die mit A auch A^* enthalten, z. B. Ringen, genügt also die Vertauschbarkeit aller Elemente⁹⁾. A. a. O. Anm. ⁹⁾ wurde übrigens gezeigt, daß M dann und nur dann Abelsch ist, wenn es der Ring $R(M)$ ist, und daß ein Ring dann und nur dann Abelsch ist, wenn er aus lauter normalen Operatoren besteht.

3. In A wurde nun bewiesen¹⁰⁾: wenn A alle normalen Operatoren durchläuft, so durchläuft $R(A)$ alle Abelschen Ringe, ja es ist zulässig, A auf die Hermiteschen zu beschränken (vgl. Anm. ⁸⁾). D. h. die Elemente des Ringes sind im weiter oben skizzierten Sinne Limites von Polynomen von A . Wir wünschen sie aber direkt als Funktionen von A darzustellen. Daher werden wir in dieser Arbeit den allgemeinen Begriff einer Funktion eines Hermiteschen (oder normalen) Operators einführen (und die für diesen Zweck zulässige Funktionenklasse genau abgrenzen), und dann beweisen: $R(A)$ ist die Menge der $f(A)$, wobei $f(x)$ alle zulässigen Funktionen mit $f(0) = 0$ durchläuft — übrigens darf man $f(x)$ auf die Funktionen der ersten Baireschen Klasse beschränken.

Da jede Abelsche Menge Teil eines Abelschen Ringes ist (z. B. dessen, den sie aufspannt), sind ihre Elemente Funktionen eines A — und jede Menge von vertauschbaren Hermiteschen Operatoren ist (da in ihr $B^* = B$ gilt) nach Definition Abelsch, so daß auch für sie das obige gilt. Wir werden die verschiedenen Formulierungs- und Verschärfungsmöglichkeiten dieses Satzes weiter unten noch näher analysieren.

4. Wie aus dem bisher Gesagten hervorgeht, ist unsere Hauptaufgabe die Einführung des allgemeinen Funktionsbegriffes für Operatoren. Hierfür sind zwei Wege vorhanden: einmal kann man es durch sukzessive Grenzübergänge für die Funktionen aller Baireschen Klassen tun, andererseits aber durch eine Definition mit Lebesgue-Stieltjesschen Integralen auf einen Schlag für die (allgemeinere) volle Funktionenklasse, für welche es überhaupt möglich ist. Die beiden Methoden sind übrigens mit zwei Wegen zur Einführung des allgemeinen Integralbegriffes, denen von Young bzw. Lebesgue¹¹⁾, aufs engste verwandt. Wir werden uns der zweiten bedienen, die allgemeinere Resultate zu erzielen gestattet, obwohl auch die erste genügen würde, um den in 3. genannten Satz zu sichern.

Ehe wir übrigens an unseren eigentlichen Gegenstand herantreten, müssen wir einen Hilfssatz über Lebesguesche Integrale beweisen, der, nebst einigen Folgerungen, vielleicht auch an und für sich nicht ganz ohne

Interesse ist und, soweit es dem Verfasser bekannt ist, in der Literatur bisher nicht vorkam.

Die Einteilung der Arbeit ist die folgende: Im Kapitel I werden die soeben erwähnten Hilfssätze über Lebesguesche Integrale bewiesen und dann das Lebesgue-Stieltjessche Integral, mitsamt seinen für unsere Zwecke wesentlichen Eigenschaften, eingeführt (da hierfür ein kurzer Exkurs in aller Strenge ausreicht, und wir es, etwas über die übliche Definitionsweise hinausgehend, auch für unstetige Differentiale brauchen, soll diese Darlegung erfolgen, obwohl sie nichts wesentlich Neues bietet). Im Kapitel II definieren wir die Funktionen von Operatoren und stellen ihre Haupteigenschaften auf; im Kapitel III wird der in 3. erwähnte Satz nebst einigen Folgerungen bewiesen. Im Kapitel IV verallgemeinern wir den Funktionsbegriff (der bis dahin nur für ein Argument, das Hermitesch sein muß, vorliegt) auf endlich oder abzählbar unendlich viele Argumente, die vertauschbare normale Operatoren sein dürfen.

I. Hilfssätze über Lebesguesche und Lebesgue-Stieltjessche Integrale.

1. Sei $\varphi(a)$ eine in einem Intervalle $0 \leq a \leq \alpha$ (α eine beliebige Zahl ≥ 0) definierte, monoton nicht-fallende, nie negative, überall nach rechts halbstetige Funktion — wir sagen kurz: eine M -Funktion.

Sei $f(x)$ eine in einer meßbaren Menge $\mathfrak{M}$ von endlichem Maße definierte, meßbare, nie negative und beschränkte Funktion. Wir definieren die Maßfunktion $\varphi(a)$ von $f(x)$ folgendermaßen:

$\varphi(a) = \text{Maß der Menge aller } x \text{ (in } \mathfrak{M}) \text{ mit } f(x) \leq a$. $\varphi(a)$ ist offenbar eine M -Funktion, als α wählen wir dabei die obere Grenze des Wertevorrats von $f(x)$.

Da jede M -Funktion $\varphi(a)$ auch die soeben genannten Eigenschaften der $f(x)$ besitzt, können wir zu $\varphi(a)$ ihre Maßfunktion, die M -Funktion $\psi(b)$ bilden. Zwischen $\varphi(a)$ und $\psi(b)$ besteht offenbar die folgende (auch zur Definition verwendbare) Relation:

$\psi(b) = \text{Obere Grenze der Menge aller } a (\geq 0, \leq \alpha) \text{ mit } \varphi(a) \leq b$.

Das Definitionsintervall von $\psi(b)$ ist $0 \leq b \leq \varphi(\alpha)$. Aus der vorangehenden Relation folgert man mühelos: Die Maßfunktion von $\psi(b)$ ist wieder $\varphi(a)$, d. h. die Beziehung zwischen $\varphi(a)$ und $\psi(b)$ ist wechselseitig. Ferner ist $\varphi(\psi(b)) = b$, wenn b nicht am linken Ende oder im Inneren eines Konstanzintervalles von ψ liegt (d. h. am unteren Ende oder im Inneren eines Intervalles, das an einer Unstetigkeitsstelle von φ durch dessen „Sprung“ ausgemacht wird); und ebenso $\psi(\varphi(a)) = a$ mit der entsprechenden Ausnahme für a . Somit sind $\varphi(a)$, $\psi(b)$ invers zueinander — soweit so etwas überhaupt möglich ist.

2. Sei $f(x)$ wie am Anfang von 1., $\varphi(a)$ ihre Maßfunktion, $\psi(b)$ die Maßfunktion von $\varphi(a)$. Wenn $f(x)$ eine M -Funktion ist, so ist, wie wir sahen, $f \equiv \psi$ (und nur dann, da $\psi(b)$ ein M ist); aber auch bei beliebigem $f(x)$ besteht eine recht innige Beziehung zwischen ihm und $\psi(b)$. Es gilt nämlich, wie wir zeigen wollen:

SATZ 1. Sei $g(u)$ eine beliebige Funktion¹²⁾. Dann gilt die Gleichung

$$\int_{\mathfrak{M}} g(f(x)) \, dx = \int_0^{\text{Maß von } \mathfrak{M}} g(\psi(b)) \, db$$

($0 \leq b \leq \text{Maß von } \mathfrak{M}$ ist das Definitionsintervall von $\psi(b)$, da das Maß von $\mathfrak{M}$ das Maximum von $\varphi(a)$ ist) in dem Sinne, daß die linke Seite immer sinnvoll ist, wenn es die rechte ist (dies kommt auf die Lebesgue-Summierbarkeit der betreffenden Integranden heraus) und beide einander dann gleich sind.

Beweis: Es genügt zu zeigen: wenn alle Mengen der b mit $g(\psi(b)) \leq m$ meßbar sind, so sind es auch alle Mengen der x mit $g(f(x)) \leq m$, und sie haben bzw. die gleichen Maße. Dies steht jedenfalls fest, wenn folgendes für alle Mengen $\mathfrak{B}$ gilt: wenn die Menge der b mit $\psi(b)$ aus $\mathfrak{B}$ meßbar ist, so ist es auch diejenige der x mit $f(x)$ aus $\mathfrak{B}$, und beide haben dasselbe Maß. (Man wähle alsdann $\mathfrak{B}$ als Menge der u mit $g(u) \leq m$.) Nennen wir die b -Menge $\mathfrak{B}_b$, die x -Menge $\mathfrak{B}_x$. Verwenden wir nunmehr $\mathfrak{B}_b$ als Ausgangspunkt. Es hat allenfalls die folgende Eigenschaft: sei I ein Konstanzintervall von $\psi(b)$ ¹³⁾, dann ist I entweder Teil von $\mathfrak{B}_b$ oder fremd zu $\mathfrak{B}_b$. Die Zahl der (paarweise fremden) Konstanzintervalle von $\psi(b)$ ist endlich oder abzählbar, sie mögen $I_1, I_2, \dots$ heißen. Wir betrachten $\mathfrak{B}_b$ nun ganz unabhängig von $\mathfrak{B}$, als beliebige Menge, nur muß jedes I_n Teil von $\mathfrak{B}_b$ oder fremd zu $\mathfrak{B}_b$ sein (und $\mathfrak{B}_b$ Teil des Intervalles $0, \text{Maß von } \mathfrak{M}$). $\mathfrak{B}_x$ ist die Menge aller x (von $\mathfrak{M}$), für die $f(x) = \psi(b)$ für ein geeignetes b von $\mathfrak{B}_b$ ist. Es gilt also zu zeigen: wenn $\mathfrak{B}_b$ meßbar ist, so ist es auch $\mathfrak{B}_x$, und beide Mengen haben dasselbe Maß.

Wir wollen zeigen: es genügt dies für diejenigen $\mathfrak{B}_b$ zu beweisen, die Intervalle $0 \leq b < \beta$ oder $0 \leq b \leq \beta$ sind (und zu den I_n im oben genannten Verhältnis stehen). Von diesen überträgt es sich nämlich zunächst auf ihre Differenzen, d. h. auf alle Intervalle (mit beliebiger Endenzurechnung, wieder das obige Verhältnis zu den I_n) — also insbesondere auf die offenen Intervalle sowie auf die I_n . Von diesen überträgt es sich wiederum auf die Summen endlich oder abzählbar vieler solcher: also erstens auf alle Summen irgendwelcher I_n , zweitens auf alle offenen Mengen (d. h. die in $0 \leq b \leq \text{Maß von } \mathfrak{M}$ relativ offenen, natürlich mit dem obigen Verhältnis zu den I_n). Insbesondere gilt sie für das volle Intervall $0 \leq b \leq \text{Maß von } \mathfrak{M}$.

Nun genügt es für alle unsere $\mathfrak{P}_b$ Maß $\mathfrak{P}_b \geq$ Äußeres Maß $\mathfrak{P}_x$ zu beweisen (ohne die Meßbarkeit von $\mathfrak{P}_x$!); denn durch Betrachtung der Komplementärmenge von $\mathfrak{P}_b$ folgt (wegen des über das volle Intervall $0 \leq b \leq$ Maß von $\mathfrak{M}$ gesagten) Maß $\mathfrak{P}_b \leq$ Inneres Maß $\mathfrak{P}_x$, also Maß $\mathfrak{P}_b =$ Äußeres Maß $\mathfrak{P}_x =$ Inneres Maß $\mathfrak{P}_x$.

Sei nun Q_b die Vereinigungsmenge aller zu $\mathfrak{P}_b$ fremden I_n , dann gilt Maß $Q_b =$ Maß Q_x , also genügt es, Maß $(Q_b + \mathfrak{P}_b) \geq$ Äußeres Maß $(Q_x + \mathfrak{P}_x)$ nachzuweisen. $Q_b + \mathfrak{P}_b$ umfaßt alle I_n , wenn wir es wieder mit $\mathfrak{P}_b$ bezeichnen, so finden wir: es genügt die obige Relation für diejenigen $\mathfrak{P}_b$ zu betrachten, die alle I_n umfassen.

Da $\mathfrak{P}_b$ Teil eines offenen $\mathfrak{P}'_b$ ist, dessen Maß das seine nur um $\leq \varepsilon$ übertrifft (dies gelingt für jedes $\varepsilon > 0$), genügt es, die alle I_n umfassenden offenen $\mathfrak{P}'_b$ zu betrachten: gilt unsere Relation für diese, so gilt sie auch für $\mathfrak{P}_b$ selbst. Für offene $\mathfrak{P}'_b$ (die obigen haben ja zu allen I_n das früher verlangte Verhältnis¹⁴⁾) gilt diese Relation aber, wie wir wissen, sogar mit dem $=$ -Zeichen.

Daher genügt es in der Tat, die Intervalle $\mathfrak{P}_b : 0 \leq b < \beta$ bzw. $0 \leq b \leq \beta$ zu betrachten. Aus dem Verhalten von $\mathfrak{P}_b$ zu den I_n folgt nun, daß es wirklich die b -Menge zu einem $\mathfrak{P}$ ist (b gehört zu $\mathfrak{P}_b$, wenn $\psi(b)$ zu $\mathfrak{P}$ gehört), und zwar ist (wegen der Monotonie von $\psi(b)$) $\mathfrak{P}$ ein Anfangsintervall des Intervalles $0, \infty$, ebenso wie $\mathfrak{P}_b$. D. h. $\mathfrak{P}$ ist eine Menge $0 \leq u < a$ oder $0 \leq u \leq a$. Da $\psi(b)$ wie $f(x)$ nie negativ sind, können wir $0 \leq u$ weglassen; da sich unsere Behauptung ($\mathfrak{P}_b, \mathfrak{P}_x$ beide meßbar und vom selben Maß) von einer aufsteigenden Folge von $\mathfrak{P}$ -Mengen sofort auf deren Vereinigungsmenge überträgt, genügt es, $u \leq a$ zu betrachten. Somit ist $\mathfrak{P}_b$ durch $\psi(b) \leq a$, $\mathfrak{P}_x$ durch $f(x) \leq a$ definiert. Daß beide Mengen meßbar sind, ist klar, und da $\psi(b), f(x)$ dieselbe Maßfunktion $\varphi(a)$ haben, haben beide dasselbe Maß: $\varphi(a)$.

Damit sind wir am Ziele.

3. Aus dem in Satz 1. formulierten Tatbestande wollen wir einige Folgerungen ziehen. Zu diesem Zweck führen wir die folgende Begriffsbildung ein: wenn die Funktion $f(x)$ (definiert für alle reellen x , mit komplexen Werten) so beschaffen ist, daß $f(\mathcal{O}(a))$ für alle monoton nichtfallenden und nach rechts halbstetigen $\mathcal{O}(a)$ eine meßbare Funktion ist, so sagen wir, es ist $\mathcal{O}$ -meßbar¹⁵⁾. Offenbar sind alle Polynome $\mathcal{O}$ -meßbar; und wenn $f_1, f_2, \dots$ eine Folge $\mathcal{O}$ -meßbarer Funktionen ist, die gegen ein f konvergieren (d. h. für jedes x gilt $f_n(x) \rightarrow f(x)$ für $n \rightarrow \infty$), so ist auch f $\mathcal{O}$ -meßbar (wegen der analogen Eigenschaft der gewöhnlichen Meßbarkeit). Somit sind alle Funktionen aller Baireschen Klassen $\mathcal{O}$ -meßbar. Ferner ist es klar, daß mit $f(x), g(x)$ auch $f(x) \pm g(x), af(x), f(x)g(x)$ $\mathcal{O}$ -meßbar sind (wieder wegen der analogen Eigenschaft der

gewöhnlichen Meßbarkeit). Es ist aber keineswegs trivial, daß folgendes gilt:

SATZ 2. Wenn $g(u)$ Φ -meßbar ist und $f(x)$ meßbar bzw. Φ -meßbar (aber reell!), so ist auch $g(f(x))$ meßbar bzw. Φ -meßbar.

Beweis: Es genügt die Meßbarkeit zu betrachten, der Fall der Φ -Meßbarkeit wird durch Ersetzen von $f(x)$ durch $f(\Phi(a))$ darauf zurückgeführt. Ferner dürfen wir $g(u)$ als beschränkt annehmen, da wir es sonst etwa durch $\text{Tghyp}(g(u))$ ersetzen können. Der Wertevorrat von $f(x)$ darf als im Intervalle $0 < u < 2$ gelegen vorausgesetzt werden, denn sonst ersetzen wir $f(x)$ durch $1 + \text{Tghyp}(f(x))$ und $g(u)$ durch $\bar{g}(u) = g(\text{Arctghyp}(u - 1))$ (für $0 < u < 2$, sonst etwa $\bar{g}(u) = 0$). Schließlich möge auch der Definitionsbereich von $f(x)$ $0 < x < 2$ sein, da wir sonst $f(\text{Arctghyp}(x - 1))$ betrachten.

Also: $g(u)$ ist beschränkt, $f(x)$ nie negativ, beschränkt, und in einer Menge von endlichem Maße definiert. Sei nun $\Phi(a)$ die Maßfunktion der Maßfunktion von $f(x)$, dann ist $g(\Phi(a))$ meßbar, also summierbar. Nach Satz 1. ist daher auch $g(f(x))$ summierbar, also meßbar.

Eine zweite, für später wichtige Tatsache ist diese:

SATZ 3. Sei $F(u)$ eine Φ -meßbare Funktion, $f_1(x), f_2(x), \dots$ irgendwelche meßbare Funktionen. Dann gibt es eine Funktion $G(u)$, die zur dritten Baireschen Klasse gehört, und für welche die Menge $\mathfrak{J}$ aller u mit $F(u) \neq G(u)$ die folgende Eigenschaft hat: die Menge aller x , für die irgendein $f_n(x)$ zu $\mathfrak{J}$ gehört, ist eine Lebesguesche Nullmenge¹⁶⁾.

Beweis: Indem wir $f_n(x)$ durch $f_n(\text{Arctghyp } x)$ ersetzen (die Abbildung $x \rightarrow \text{Arctghyp } x$ bildet $-1 < x < 1$ auf alle x ab, und führt dabei Lebesguesche Nullmengen in ebensolche über), erreichen wir, daß sein Definitionsbereich $-1 < x < 1$ wird; und indem wir es sodann durch $f_n\left(\frac{2}{\beta - \alpha}\left(x - \frac{\alpha + \beta}{2}\right)\right)$ ersetzen ($\alpha < \beta$, sonst beide beliebig), machen wir $\alpha < x < \beta$ zum Definitionsbereich. Dies sei daher von vornherein vorausgesetzt, und zwar habe $f_n(x)$ den Definitionsbereich $\frac{1}{n+1} < x < \frac{1}{n}$. Alle diese Funktionen

$f_n(x)$ ($n = 1, 2, \dots, \frac{1}{n+1} < x < \frac{1}{n}$) können wir nun zu einer einzigen $\bar{f}(x)$ zusammenfassen, die dann in $0 < x < 1$ definiert ist (für $x = \frac{1}{2}, \frac{1}{3}, \dots$ sei sie etwa $= 0$) und meßbar. Es handelt sich also darum, $G(u)$ so zu wählen, daß die x mit $\bar{f}(x)$ aus $\mathfrak{J}$ eine Lebesguesche 0-Menge bilden. Sei nun $\psi(b)$ die Maßfunktion der Maßfunktion von $\bar{f}(x)$, nach Satz 1. ist die x -Menge der $\bar{f}(x)$ aus $\mathfrak{J}$ gewiß eine Lebesguesche Nullmenge, wenn die b -Menge der $\psi(b)$ aus $\mathfrak{J}$ eine ist (man wähle $g(u) = 1$ für die u von $\mathfrak{J}$, sonst $= 0$). Also dürfen wir $f(x)$ durch $\psi(b)$ ersetzen.

D. h.: es genügt den Fall zu betrachten, wo nur ein einziges $f_1(x)$ auftritt, und dieses eine M -Funktion ist (vgl. Anm. ¹⁶). Wir nennen es $\varphi(x)$ (statt $\psi(b)$).

Zunächst ist $F(u)$ Φ -meßbar, also $F(\varphi(x))$ meßbar. Auf Grund des Vitalischen Satzes (vgl. Anm. ¹⁶) existiert also eine Funktion $\bar{G}(x)$ aus der zweiten Baireschen Klasse, so daß bis auf eine Lebesguesche Nullmenge immer $F(\varphi(x)) = \bar{G}(x)$ gilt. Hätte nun $\bar{G}(x)$ die Form $G(\varphi(x))$, mit einem $G(u)$ aus der dritten Baireschen Klasse, so wäre damit der Beweis fertig. Er ist es aber auch, wenn es uns gelingt, $\bar{G}(x)$ durch Abänderungen in einer Lebesgueschen Nullmenge in ein solches $\bar{G}(x) = G(\varphi(x))$ zu verwandeln. Dies soll daher geschehen.

Daß $\bar{G}(x)$ die Form $G(\varphi(x))$ (für irgendein $G(u)$) hat, besagt bloß, daß es in allen Konstanzintervallen von $\varphi(x)$ konstant ist. $F(\varphi(x))$ ist so, daher ist es $\bar{G}(x)$ bis auf eine Lebesguesche Nullmenge. Indem wir also $\bar{G}(x)$ in den Konstanzintervallen von $\varphi(x)$ konstant machen, verändern wir es bloß in einer Lebesgueschen Nullmenge. So entsteht ein $\bar{\bar{G}}(x) = G(\varphi(x))$; es ist nur noch zu zeigen, daß $G(u)$ von der dritten Baireschen Klasse ist. Übrigens können wir, wenn $\psi(u)$ die Maßfunktion von $\varphi(x)$ ist, $G(u) = \bar{G}(\psi(u))$ setzen, da $\psi(\varphi(x)) = x$ ist, falls x nicht in einem Konstanzintervall von $\varphi(x)$ liegt, und wenn es in einem solchen liegt, so liegen $\psi(\varphi(x))$ und x im nämlichen, und $\bar{G}(x)$ ist dort konstant. Also ist $\bar{\bar{G}}(x) = G(\varphi(x))$.

Zunächst ist $\bar{\bar{G}}(x)$ von der zweiten Baireschen Klasse: denn es entsteht aus einer solchen Funktion, $\bar{G}(x)$, indem diese in endlich oder abzählbar vielen Intervallen (den Konstanzintervallen von $\varphi(x)$) abgeändert, und zwar konstant gemacht, wird. Dies ist offenbar für die zweite Bairesche Klasse allgemein richtig, wenn es bei einer endlichen Zahl von Intervallen für die erste Bairesche Klasse stimmt¹⁷). Eine Funktion der ersten Baireschen Klasse ist nun ein Limes einer überall konvergenten Folge stetiger Funktionen, wird sie auf die genannte Weise geändert, so ist sie Limes der Folge der ebenso abgeänderten stetigen Funktionen. Diese sind aber auch nach dieser Änderung alle überall stetig — mit Ausnahme endlich vieler Stellen (der Enden der Änderungsintervalle). Man erkennt leicht, daß die Limes solcher Funktionen auch als Limes stetiger Funktionen darstellbar sind¹⁸). Als Letztes müssen wir noch zeigen: da $\bar{\bar{G}}(x)$ zur zweiten Baireschen Klasse gehört, gehört $G(u) = \bar{G}(\psi(u))$ zur dritten. Nun entsteht $\bar{\bar{G}}(x)$ durch zwei sukzessive Grenzübergänge aus gewissen stetigen Funktionen $f_{m,n}(x)$ ¹⁹), also $G(u) = \bar{G}(\psi(u))$ ebenso aus den $f_{m,n}(\psi(u))$, welche (wie $\psi(u)$) nach rechts halbstetig sind. Da jede nach rechts halbstetige Funktion zur ersten Baireschen Klasse gehört²⁰), gehört $G(u)$ zur dritten.

Damit schließen wir unsere aufs Lebesguesche Maß und Integral bezüglichen Betrachtungen ab.

4. Das Lebesgue-Stieltjessche Integral $\int f(x) d\varphi(x)$ soll nun definiert werden, und zwar zunächst für den Fall, wo $\varphi(x)$ eine M -Funktion ist. Dann ist der Definitionsbereich von $\varphi(x)$ ein Intervall $0 \leq x \leq \alpha$, daher sollen auch die Integrationsgrenzen diese Zahlen sein — wir halten sie zunächst fest. Die Maßfunktion von $\varphi(x)$ sei $\psi(y)$, dann definieren wir (an Lebesgue anlehnd):

$$\int f(x) d\varphi(x) = \int f(\psi(y)) dy,$$

wo die rechte Seite ein gewöhnliches Lebesguesches Integral ist. Sie ist sicher sinnvoll, wenn $f(x)$ Φ -meßbar und beschränkt ist — darum wollen wir dies im folgenden stets über $f(x)$ voraussetzen.

Das so definierte Integral hat die folgenden Eigenschaften:

$$a) \int a f(x) \cdot d\varphi(x) = a \int f(x) \cdot d\varphi(x) \quad (a \text{ eine komplexe Zahl}).$$

$$b) \int (f(x) + g(x)) \cdot d\varphi(x) = \int f(x) \cdot d\varphi(x) + \int g(x) \cdot d\varphi(x).$$

$$c) \int f(x) \cdot da\varphi(x) = a \int f(x) \cdot d\varphi(x) \quad (a > 0).$$

$$d) \int f(x) \cdot d(\varphi(x) + \chi(x)) = \int f(x) \cdot d\varphi(x) + \int f(x) \cdot d\chi(x).$$

e) Gelte $f(x) \geq 0$ in einer Menge $\mathfrak{M}$, wobei die Menge der y für die $\psi(y)$ nicht zu $\mathfrak{M}$ gehört, eine Lebesguesche Nullmenge ist²¹⁾. Dann ist $\int f(x) \cdot d\varphi(x) \geq 0$.

f) Sei $f(x) = g(x)$, in einer Menge $\mathfrak{M}$ wie in e), dann ist

$$\int f(x) \cdot d\varphi(x) = \int g(x) \cdot d\psi(x).$$

g) Sei $f_n(x) \rightarrow f(x)$ für $n \rightarrow \infty$, in einer Menge $\mathfrak{M}$ wie in e), ferner seien die $f_n(x)$ gleichmäßig beschränkt. Dann gilt

$$\int f_n(x) \cdot d\varphi(x) \rightarrow \int f(x) \cdot d\varphi(x).$$

Da a), b) sofort aus den entsprechenden Eigenschaften des Lebesgueschen Integrals folgen, und auch e), f), g) durch die Definition sofort in bekannte Eigenschaften desselben übergehen, sind nur c), d) zu prüfen. Aber nach f) und nach SATZ 3 genügt es, diese für $f(x)$, $g(x)$ aus der dritten Baireschen Klasse zu beweisen, und wegen g) sogar für die der nullten, d. h. für stetige $f(x)$, $g(x)$. Dann aber haben wir es offenbar mit gewöhnlichen

Riemann-Stieltjesschen Integralen zu tun (denn $\int f(\psi(y)) dy$ ist, da $f(\psi(y))$ nach rechts halbstetig ist, ein Riemannsches Integral), und für diese gelten c), d) bekanntlich.

Ferner erkennt man: wenn eine monoton wachsende stetige, also ein-eindeutige, Abbildung $\xi = \omega(x)$, $x = \omega^{-1}(\xi)$ von $0 \leq x \leq \alpha$ auf $0 \leq \xi \leq \beta$ gegeben ist, so ist

$$\int f(\xi) d\varphi(\xi) = \int f(\omega(x)) d\varphi(\omega(x)).$$

(Daß $f(\omega(x))$ $\mathcal{O}$ -meßbar ist, folgt offenbar aus der $\mathcal{O}$ -Meßbarkeit von $f(\xi)$). Es ist nämlich, wie man sofort erkennt, $\omega^{-1}(\psi(y))$ die Maßfunktion von $\varphi(\omega(x))$ ($\psi(y)$ sei diejenige von $\varphi(x)$), also $f(\omega[\omega^{-1}(\psi(y))]) = f(\psi(y))$.

Dies ermöglicht uns nun, die Definition von $\int f(x) d\varphi(x)$ auf den Fall zu übertragen, wo $\varphi(x)$ für alle reellen Zahlen definiert ist, und entsprechend von $-\infty$ bis $+\infty$ integriert wird (im übrigen sei wieder $f(x)$ $\mathcal{O}$ -meßbar und beschränkt, $\varphi(x)$ monoton nichtfallend, nach rechts halbstetig, und beschränkt). Sei nämlich $\xi = \varrho(x)$, $x = \varrho^{-1}(\xi)$ eine monoton wachsende stetige, also ein-eindeutige, Abbildung von $0 < x < \alpha$ auf $-\infty < \xi < \infty$. Wir setzen dann

$$\int f(\xi) d\varphi(\xi) = \int f(\varrho(x)) d\varphi(\varrho(x)).$$

Von der Wahl von $\varrho(x)$ ist die definierende rechte Seite unabhängig: um dies für zwei $\varrho_1(x)$, $\varrho_2(x)$ zu erkennen, genügt es, im obigen Resultat $\omega(x) = \varrho_2(\varrho_1^{-1}(x))$ zu setzen. Auch für diesen Integralbegriff gelten a) — g), wie man ohne weiteres sieht. (Bei e) — g) ist die Forderung über $\mathfrak{M}$ in der Formulierung in Anm. ²¹⁾ zu stellen, da wir die „Maßfunktion“ für das allgemeine $\varphi(x)$ nicht definiert haben.) Ferner sind auch bei diesem $\int$ die Variablentransformationen $\xi = \varrho(x)$, $x = \varrho^{-1}(\xi)$ unwirksam, wie aus der Definition folgt, nur müssen alle x auf alle ξ abgebildet werden. Mit Hilfe von $\int$, genauer $\int_{-\infty}^{\infty}$, können wir nun auch alle $\int_{\mathfrak{M}}$ definieren ²²⁾: wenn $f(x)$ gegeben ist, so sei $f_{\mathfrak{M}}(x) = f(x)$ in $\mathfrak{M}$ und $= 0$ in der Komplementärmenge, und

$$\int_{\mathfrak{M}} f(x) \cdot d\varphi(x) = \int f_{\mathfrak{M}}(x) \cdot d\varphi(x).$$

Aus der Gültigkeit von a) — g) für $\int$ folgt sofort diejenige für $\int_{\mathfrak{M}}$, ferner folgt aus der Definition von $\int_{\mathfrak{M}}$ und aus b), g) sofort:

h) Seien $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ paarweise elementfremde $\mathcal{O}$ -meßbare Mengen (in endlicher oder abzählbarer Zahl), dann ist ($\overline{\mathfrak{M}} = \mathfrak{M}_1 + \mathfrak{M}_2 + \dots$):

$$\int_{\mathfrak{M}_1} f(x) \cdot d\varphi(x) + \int_{\mathfrak{M}_2} f(x) \cdot d\varphi(x) + \dots = \int_{\mathfrak{M}} f(x) \cdot d\varphi(x).$$

Schließlich ist $\int_{\mathfrak{M}}$ mit dem eingangs definierten $\int$ im Intervalle $0, \alpha$ identisch, wenn $\mathfrak{M}$ die Menge $0 < x < \alpha$ ist. Sei nämlich $f(x)$ außerhalb von $0 < x < \alpha$ als 0 definiert und $\varphi(x)$ als $\varphi(0)$ bzw. als $\varphi(\alpha)$, und betrachten wir $\int f(x-\varepsilon) d\varphi(x-\varepsilon)$ für $0 < x < \alpha + 2\varepsilon$ (ε irgendeine feste Zahl > 0). Sei $\psi(y)$ die Maßfunktion von $\varphi(x)$ (in $0, \alpha$), dann ist die Maßfunktion von $\varphi(x+\varepsilon)$ (in $0, \alpha + 2\varepsilon$) gleich $\psi(y) + \varepsilon$ für $0 \leq y < \varphi(\alpha)$ und $= \alpha + 2\varepsilon$ (nicht, wie $\psi(y) + \varepsilon, = \alpha + \varepsilon$) für $y = \varphi(\alpha)$ ($0 \leq \alpha \leq \varphi(\alpha)$ ist ja der Definitionsbereich). Hieraus folgt $\int_{0, \alpha + 2\varepsilon} f(x-\varepsilon) \cdot d\varphi(x-\varepsilon) = \int_{0, \alpha} f(x) d\varphi(x)$. Nun ist aber die linke Seite gleich $\int f(x) d\varphi(x)$ (von $-\infty$ bis $+\infty$), es genügt dazu das $0, \alpha + 2\varepsilon$ auf $-\infty, +\infty$ abbildende $\varrho(x)$ in $\varepsilon < x < \alpha + \varepsilon$ gleich $x - \varepsilon$ zu wählen, was möglich ist. Und dieses Integral ist, da $f(x)$ außerhalb $0 < x < \alpha$ verschwindet, gleich $\int_{\mathfrak{M}} f(x) \cdot d\varphi(x)$, wenn $\mathfrak{M}$ diese Menge ist.

Es ist noch erwähnenswert, daß die Variablentransformationsformel auf Grund der Definition die folgende Form annimmt:

$$\int_{\mathfrak{M}} f(\xi) \cdot d\varphi(\xi) = \int_{\varrho^{-1}(\mathfrak{M})} f(\varrho(x)) \cdot d\varphi(\varrho(x)),$$

wo $\varrho^{-1}(\mathfrak{M})$ das $x = \varrho^{-1}(\xi)$ -Bild von $\mathfrak{M}$ ist. Ferner können wir in e) — g) (Anm. ²¹) statt der Komplementärmenge von $\mathfrak{M}$ den in $\mathfrak{M}$ liegenden Teil derselben nehmen: denn außerhalb von $\mathfrak{M}$ ist ja ohnehin $f_{\mathfrak{M}}(x) = g_{\mathfrak{M}}(x) = f_{n, \mathfrak{M}}(x) = 0$, also alles erfüllt.

5. Wir gehen nun dazu über, $\int_{\mathfrak{M}} f(x) d\varphi(x)$ für solche $\varphi(x)$ zu definieren, die (im ganzen Intervall $-\infty, +\infty$) von beschränkter Schwankung sind. ($f(x)$, $\mathfrak{M}$ seien wie früher: $\mathcal{O}$ -meßbar und beschränkt, bzw. $\mathcal{O}$ -meßbar.) Ein solches $\varphi(x)$ kann bekanntlich immer auf diese Form gebracht werden:

$$\varphi(x) = \psi_1(x) - \psi_2(x) + i(\psi_3(x) - \psi_4(x)),$$

wo $\psi_1, \psi_2, \psi_3, \psi_4$ M -Funktionen (d. h. reell, nie negativ [$\varphi(x)$ darf komplex sein!], monoton nichtfallend und nach rechts halbstetig) sind —

d. h. diese Gleichung gilt mit Ausnahme der (endlich oder abzählbar vielen) Stellen x , wo $\varphi(x)$ nicht nach rechts halbstetig ist²⁹⁾. Wir definieren nun:

$$\int_{\mathfrak{M}} f(x) d\varphi(x) = \int_{\mathfrak{M}} f(x) d\psi_1(x) - \int_{\mathfrak{M}} f(x) d\psi_2(x) + \\ + i \left(\int_{\mathfrak{M}} f(x) d\psi_3(x) - \int_{\mathfrak{M}} f(x) d\psi_4(x) \right).$$

Daß dies von der speziellen Wahl von $\psi_1, \psi_2, \psi_3, \psi_4$ unabhängig ist, erkennt man, indem man beachtet, daß $\int_{\mathfrak{M}} f(x) d\chi_1(x) - \int_{\mathfrak{M}} f(x) d\chi_2(x)$ (χ_1, χ_2 M -Funktionen) nur von $\chi_1(x) - \chi_2(x)$ abhängt, was unmittelbar aus d) folgt. Daß es für M -Funktionen $\varphi(x)$ die alte Definition ist, folgt daraus, daß dann $\psi_1 \equiv \varphi, \psi_2 \equiv \psi_3 \equiv \psi_4 \equiv 0$ gesetzt werden kann.

a)–h) können auch auf diesen Integralbegriff übertragen werden, wie man aus der Definition sofort abliest, ebenso die Variablentransformationsformel am Ende von 4. Immerhin ist noch hinzuzufügen: Ad c) Dies gilt für alle komplexen a . Es genügt ja, es für M -Funktionen $\varphi(x)$ zu beweisen, und bei diesen, da es für positive a ohnehin gilt (und wegen d)) für $a = \pm 1, \pm i$. Dies sind Potenzen von i , also genügt $a = i$, und das liest man an der Definition ab. Ad e)–g), die in Anm. ²¹⁾ für $\mathfrak{M}$ formulierte Bedingung. Hier sind die $\psi_1, \psi_2, \psi_3, \psi_4$ einzusetzen, und nicht etwa φ selbst! Immerhin ist zu beachten, daß bei der Wahl der $\psi_1, \psi_2, \psi_3, \psi_4$ nach Anm. ²³⁾ diese nur dort unstetig sind, wo φ es ist. Ad e) Mit $\geq 0, \leq 0$ gilt dies nur für M -Funktionen $\varphi(x)$, mit $= 0$ aber offenbar allgemein. Wenn $f(x), \varphi(x)$ reell sind, ist das Integral offenbar auch reell.

Zum Schluß beweisen wir noch die folgende wichtige Relation:

i) Seien $f(x), g(x)$ $\mathcal{O}$ -meßbar und beschränkt, $\varphi(x)$ von beschränkter Schwankung. Dann ist auch $\int_{\mathfrak{M}}^x g(y) \cdot d\varphi(y)$ (als Funktion von x betrachtet, $\int_{\mathfrak{M}}^x$ sei das Integral über die Menge der $y \leq x$) von beschränkter Schwankung, und es gilt:

$$\int_{\mathfrak{M}} f(x) \cdot d \left[\int_{\mathfrak{M}}^x g(y) \cdot d\varphi(y) \right] = \int_{\mathfrak{M}} f(x) g(x) \cdot d\varphi(x).$$

Da sich jedes $f(x)$ und jedes $g(x)$ mit den Faktoren $\pm 1, \pm i$ additiv aus nichtnegativen (und dabei $\mathcal{O}$ -meßbaren und beschränkten) zusammensetzen läßt, und ebenso jedes $\varphi(x)$ aus nichtnegativen monoton nichtfallenden und beschränkten (wegen a)–d)), dürfen wir uns von vornherein

auf solche beschränken. Dann ist aber $\int_x^x g(y) \cdot d\varphi(y)$ monoton nichtfallend und beschränkt ($\geq 0, \leq \int g(y) \cdot d\varphi(y)$ nach e), h)), also von beschränkter Schwankung — d. h. die erste Behauptung bewiesen. Es bleibt noch der Beweis der Integralformel zu führen. Nach f) und Satz 3 brauchen wir nur solche $f(x), g(x)$ zu betrachten, die zur dritten Baireschen Klasse gehören, und wegen der in g) zum Ausdruck kommenden „Stetigkeit“ beider Seiten in $f(x)$, können wir $f(x)$ sogar auf die nullte Bairesche Klasse beschränken: d. h. es als stetig annehmen. Aber für stetiges $f(x)$ ist die linke Seite, $\int f(x) d\eta(x)$ ($\eta(x) = \int_x^x g(y) \cdot d\varphi(y)$), ein gewöhnliches Riemann-Stieltjessches Integral, also in $\eta(x)$ „stetig“. Und nach g) sind sowohl $\eta(x)$, als auch die rechte Seite, $\int f(x) g(x) \cdot d\varphi(x)$, in $g(x)$ „stetig“, so daß wir auch $g(x)$ als zur nullten Baireschen Klasse gehörig, d. h. stetig, ansehen dürfen. Dann aber haben wir es auf beiden Seiten nur noch mit Riemann-Stieltjesschen Integralen zu tun, so daß beide Seiten auch in $\varphi(x)$ „stetig“ sind. Und da es sicher eine überall gegen $\varphi(x)$ konvergierende Folge überall stetiger und differentiierbarer (nichtnegativer monoton nichtfallender) Funktionen gibt, können wir $\varphi(x)$ als solche voraussetzen. Dann können wir die Riemann-Stieltjesschen Integrale als Riemannsche Integrale schreiben:

$$\begin{aligned} \int f(x) \cdot d \left[\int_x^x g(y) \cdot d\varphi(y) \right] &= \int f(x) \cdot \frac{d}{dx} \left[\int_x^x g(y) \cdot \frac{d}{dy} \varphi(y) \cdot dy \right] \cdot dx \\ &= \int f(x) g(x) \cdot \frac{d}{dx} \varphi(x) \cdot dx, \\ \int f(x) g(x) \cdot d\varphi(x) &= \int f(x) g(x) \cdot \frac{d}{dx} \varphi(x) \cdot dx. \end{aligned}$$

Aber dies setzt die Gültigkeit der behaupteten Relation in Evidenz.

II. Allgemeine Funktionen von Operatoren.

1. Sei A ein beschränkter Hermischer Operator, somit einer aus $\mathcal{E}$. Wir schreiben A in der Hilbertschen Spektralform²⁴⁾:

$$(A f, g) = \int \lambda d(E(\lambda) f, g)$$

(λ ist die Integrationsvariable), wo $E(\lambda)$ die zu A gehörige Zerlegung der Einheit ist. (Die Bezeichnungsweise ist die in E und A , a. a. O. Anm. ²⁴⁾), auseinandergesetzte. Unter $f, g, \dots, \varphi, \psi, \dots$ verstehen wir von nun an

stets Funktionen des Hilbertschen Raumes, allgemeine [im Reellen definierte] Funktionen nennen wir $F, G, \dots$). Das Integral ist ein Riemann-Stieltjessches²⁵⁾, und gewiß (trotz des unbeschränkten λ !) sinnvoll, da $E(\lambda)$ (also auch $(E(\lambda)f, g)$) außerhalb eines endlichen Intervalles konstant ist²⁶⁾, so daß es genügt, über dieses zu integrieren.

Betrachten wir $E(\lambda)$ als Funktion von λ . Im Sinne der für Projektionsoperatoren definierten Relation $\leq$ ist es monoton nichtfallend, nach rechts halbstetig²⁷⁾, und unterhalb einer geeigneten Schranke $= 0$, oberhalb einer geeigneten Schranke $= 1$. Die Vereinigungsmenge seiner (offen genommenen) Konstanzintervalle heiße $\mathfrak{I}$, die Komplementärmenge von $\mathfrak{I}$ heiße $\mathfrak{S}$. $\mathfrak{I}$ ist offen, $\mathfrak{S}$ abgeschlossen, $\mathfrak{S}$ ist nach dem vorhin Gesagten beschränkt. Für alle f, g ist $(E(\lambda)f, g)$ in den Konstanzintervallen von $E(\lambda)$ konstant, d. h. $\mathfrak{I}$ durch die Konstanzintervalle dieser Funktion überdeckt. $\mathfrak{S}$ ist bekanntlich das „Spektrum“ von A . Ferner sei $\mathfrak{P}$ die Menge aller Unstetigkeitsstellen von $E(\lambda)$. Nur an diesen kann ein $(E(\lambda)f, g)$ unstetig sein (vgl. Anm. 27)). $\mathfrak{P}$ ist natürlich Teilmenge von $\mathfrak{S}$ und bekanntlich endlich oder abzählbar, es ist das „Punktspektrum“ von A .

Sei nun $F(x)$ eine beschränkte $\mathcal{O}$ -meßbare Funktion. Wir betrachten die Ausdrücke

$$\int F(\lambda) d(E(\lambda)f, g).$$

Da $\mathfrak{I}$ aus lauter Konstanzintervallen der Funktion hinter dem d -Zeichen besteht, sind diese Ausdrücke alle vom Verhalten von $F(x)$ in $\mathfrak{I}$ unabhängig (I. 4. und 5., Eigenschaft f), zusammen mit Anm. 21)). Daher können wir z. B. $F(x)$ in $\mathfrak{I} = 0$ setzen (da $\mathfrak{I}$ offen ist, bleibt $F(x)$ dabei $\mathcal{O}$ -meßbar), oder, was dasselbe ist, $\int$ durch $\int_{\mathfrak{S}}$ ersetzen. Infolgedessen braucht $F(x)$ gar nicht beschränkt zu sein es genügt, wenn es in $\mathfrak{S}$ beschränkt ist²⁸⁾.

Es gilt allgemein die folgende Abschätzung

$$\left| \int F(\lambda) d(E(\lambda)f, g) \right| \leq \sqrt{\int |F(\lambda)| \cdot d(|E(\lambda)f|^2) \cdot \int |F(\lambda)| \cdot d(|E(\lambda)g|^2)}.$$

Zunächst genügt es, dieselbe für die $F(\lambda)$ der dritten Baireschen Klasse zu beweisen (Satz 3 und f)), und infolgedessen sogar für diejenigen der nullten (g)), d. h. für die stetigen. Dann handelt es sich um Riemann-Stieltjessche Integrale, die somit bzw. Limites von

$$\begin{aligned} & \sum_1^N F(\lambda_n) \{(E(\lambda_n)f, g) - (E(\lambda_{n-1})f, g)\} \\ &= \sum_1^N F(x_n) \{(E(\lambda_n) - E(\lambda_{n-1}))f, g\}, \end{aligned}$$

$$\begin{aligned}
& \sum_1^N {}^n |F(\lambda_n)| \{ |E(\lambda_n) f|^2 - |E(\lambda_{n-1}) f|^2 \} \\
&= \sum_1^N {}^n |F(x_n)| \{ |E(\lambda_n) - E(\lambda_{n-1})\} f|^2 {}^{29}), \\
& \sum_1^N {}^n |F(\lambda_n)| \{ |E(\lambda_n) g|^2 - |E(\lambda_{n-1}) g|^2 \} \\
&= \sum_1^N {}^n |F(x_n)| \{ |E(\lambda_n) - E(\lambda_{n-1})\} g|^2 {}^{29})
\end{aligned}$$

$(\lambda_0 < \lambda_1 < \dots < \lambda_N, \quad \lambda_0 \rightarrow -\infty, \quad \lambda_N \rightarrow +\infty, \quad \text{Max}_{n=1, \dots, N} (\lambda_n - \lambda_{n-1}) \rightarrow 0)$
sind. Wegen

$$|\{E(\lambda_n) - E(\lambda_{n-1})\} f, g| \leq |\{E(\lambda_n) - E(\lambda_{n-1})\} f| \cdot |\{E(\lambda_n) - E(\lambda_{n-1})\} g| {}^{30}),$$

und der Schwarzschen Ungleichheit

$$\left| \sum_1^N {}^n a_n b_n \right| \leq \sqrt{\sum_1^N {}^n |a_n|^2 \sum_1^N {}^n |b_n|^2}$$

ist aber der erste Ausdruck absolut genommen $\leq$ als die Quadratwurzel des Produkts der beiden letzten, und daraus folgt die Behauptung.

Gilt nun allgemein $|F(x)| \leq C$ (wir wissen übrigens, daß es genügt, wenn dies in $\mathfrak{S}$ gilt), so ist

$$\begin{aligned}
\int |F(x)| \cdot d|E(\lambda) f|^2 &\leq \int C \cdot d|E(\lambda) f|^2 = C \cdot (|1 f|^2 - |0 f|^2) = C \cdot |f|^2, \\
\int |F(x)| \cdot d|E(\lambda) g|^2 &\leq \int C \cdot d|E(\lambda) g|^2 = C \cdot (|1 g|^2 - |0 g|^2) = C \cdot |g|^2,
\end{aligned}$$

also

$$\left| \int F(x) \cdot d(E(\lambda) f, g) \right| \leq C \cdot |f| \cdot |g|.$$

Wenn wir bei festgehaltenem f $\int F(x) \cdot d(E(\lambda) f, g) = L(g)$ als Funktion von g ansehen, so haben wir darum:

$$L(a_1 g_1 + \dots + a_n g_n) = \bar{a}_1 L(g_1) + \dots + \bar{a}_n L(g_n), \quad |L(g)| \leq D \cdot |g|$$

($D = C \cdot |f|$).

Nach einem Satze von F. Rieß³¹⁾ existiert daher ein und nur ein f^* , so daß stets $(f^*, g) = L(g)$ gilt. Wir nennen es $A'f$ (es hängt ja von f ab), so daß

$$(A'f, g) = \int F(\lambda) \cdot d(E(\lambda) f, g)$$

definitiv ist.

Wie man sofort erkennt, ist A' ein linearer Operator, und $|(A'f, g)| \leq C \cdot |f| \cdot |g|$ hat die Folge $|A'f| \leq C \cdot |f|$ (man setze $g = A'f$), so daß er auch beschränkt ist — d. h. A' gehört zu **B**. Diesen Operator A' wollen wir nun $F(A)$ nennen.

2. Wir zählen die einfachsten Eigenschaften von $F(A)$ auf:

- a) Für $F(x) \equiv 0$ bzw. $\equiv 1$ bzw. $\equiv x$ ist $F(A) = 0$ bzw. $= 1$ bzw. $= A$.
- b) Für $F(x) \equiv \overline{G(x)}$ (konjugiert-komplex) ist $F(A) = G(A)^*$ ³²⁾.
- c) Für $F(x) \equiv aG(x)$ ist $F(A) = aG(A)$.
- d) Für $F(x) \equiv G(x) + H(x)$ ist $F(A) = G(A) + H(A)$.
- e) Für $F(x) \equiv G(x) \cdot H(x)$ ist $F(A) = G(A) \cdot H(A)$.
- f) Für $F(x) \equiv G(H(x))$ ist $F(A) = G(H(A))$ ³³⁾.
- g) Die Gültigkeit von $F(x) = G(x)$ in einer Menge $\mathfrak{M}$ zieht $F(A) = G(A)$ nach sich, wenn $\mathfrak{M}$ die folgende Eigenschaft hat: Sei $\xi = \varrho(x)$, $x = \varrho^{-1}(\xi)$ eine monoton wachsende stetige, d. h. ein-eindeutige, Abbildung eines Intervalles $0 < x < \alpha$ auf $-\infty < \xi < \infty$. Wie $\varrho(x)$ gewählt wird, ist gleichgültig, denn das gleich zu definierende $\varrho(\psi_f(y))$ hängt, wie man leicht erkennt, nur scheinbar davon ab. $\varrho(\psi_f(y))$ ist nämlich im am Ende von I., 1. skizzierten Sinne aufzufassende Inverse von $|E(\xi)f|^2$. Wir bilden $\Phi_f(x) = |E(\varrho(x))f|^2$ und seine Maßfunktion $\psi_f(y)$, ferner $\varrho(\psi_f(y))$. Für jedes f muß nun die Menge aller y , für die $\varrho(\psi_f(y))$ nicht zu $\mathfrak{M}$ gehört, eine Lebesguesche Nullmenge sein ³⁴⁾.
- h) Gelte $F_n(x) \rightarrow G(x)$ für $n \rightarrow \infty$, in einer Menge $\mathfrak{M}$ wie in g), ferner seien die $F_n(x)$ (wenigstens in einer Menge $\mathfrak{M}$ wie vorhin) gleichmäßig beschränkt. Dann gilt $F_n(A) \rightarrow G(A)$, und zwar im Sinne der „starken Doppelkonvergenz“ in **B** ³⁵⁾, d. h. es ist $F_n(A)f \rightarrow G(A)f$, $F_n(A)^*f \rightarrow G(A)^*f$ („stark“ im Hilbertschen Raume).

Die Behauptungen a)–h) gehen dahin, daß sich eine Reihe von Eigenschaften von den reellen Zahlen x (für die allein unsere Funktionen unmittelbar definiert sind), auf alle Hermiteschen beschränkten Operatoren übertragen. Wir wollen sie nun beweisen.

a)–d) und g) folgen sofort aus der Definition der $F(A)$, ... (d. h. der $(F(A)f, g)$, ...). Dabei ist bei g) zu beachten, daß die dortige Fassung der Bedingung für $\mathfrak{M}$ darum die richtige ist, weil die in den betreffenden Definitionen hinter dem d -Zeichen stehenden Funktionen von beschränkter Schwankung, $(E(\lambda)f, g)$, als Linearaggregate von vier $|E(\lambda)h|^2$ (mit $h = \frac{f \pm g}{2}$,

$\frac{f \pm ig}{2}$ und den bzw. Koeffizienten $\pm 1, \pm i$) aufzufassen sind (vgl. Anm. ²⁵⁾).

Und da die Integrale über das Intervall $-\infty < \lambda < \infty$ erstreckt sind, müssen sie dann noch durch $\lambda = \varrho(x)$, $x = \varrho^{-1}(\lambda)$ auf ein Intervall $0 < x < \alpha$ reduziert werden (vgl. I., 4.), was genau zu unserer $\mathfrak{M}$ -Bedingung führt.

Noch eines ist erwähnenswert: Da im zu beweisenden $(F(A)f, g) = (G(A)f, g)$ beide Seiten in f stetig sind, genügt es, dies für f, g aus einer überall dichten Menge zu beweisen, welche (da der Hilbertsche Raum so parabel ist³⁶⁾) abzählbar gewählt werden kann. Dann bilden auch die für h (in g) heißt es f) in Frage kommenden Werte $\left(\frac{f \pm g}{2}, \frac{f \pm ig}{2}\right)$ eine abzählbare Menge. Also: es genügt, die Forderungen für $\mathfrak{M}$ in g) für eine geeignete abzählbare Menge von f zu stellen.

Ebenso von selbst ergibt sich die Richtigkeit von h), wenn darin $F_n(A) \rightarrow F(A)$ bloß im Sinne der „schwachen“ Konvergenz in $\mathfrak{B}$ ausgesprochen wird (vgl. a. a. O. Anm. ³⁵⁾), d. h. wenn nur $(F_n(A)f, g) \rightarrow (F(A)f, g)$ für $n \rightarrow \infty$ zu beweisen ist. Um auch die vorige Zusatzbemerkung über g) zu übertragen (d. i. die Hinreichendheit des Betrachtens einer geeigneten abzählbaren Menge von f in der Bedingung für $\mathfrak{M}$), genügt es, sich zu vergegenwärtigen, daß die $F_n(x)$ ($n = 1, 2, \dots$) gleichmäßig beschränkte Funktionen von x sind, also die $F_n(A)$ gleichmäßig beschränkte Operatoren (vgl. 1.), also die $(F_n(A)f, g)$ gleichmäßig stetige Funktionen von f, g .

In e) ist für alle f, g die Gleichheit von

$$(F(A)f, g) = \int F(\lambda) d(E(\lambda)f, g) = \int G(\lambda) H(\lambda) d(E(\lambda)f, g)$$

und von

$$\begin{aligned} (G(A)H(A)f, g) &= (H(A)f, (G(A))^*g) = \int H(\lambda) d(E(\lambda)f, (G(A))^*g) \\ &= \int H(\lambda) d(G(A)E(\lambda)f, g) \end{aligned}$$

zu beweisen. Nun ist

$$\begin{aligned} (G(A)E(\lambda)f, g) &= \int G(\lambda') d(E(\lambda')E(\lambda)f, g) \\ &= \int G(\lambda') d(E(\text{Min}(\lambda', \lambda)f, g)). \end{aligned}$$

Der Ausdruck hinter dem d -Zeichen ist für $\lambda' \geq \lambda$ konstant ($\text{Min}(\lambda', \lambda) = \lambda$, λ ist fest, λ' die Variable!), so daß wir das Integral durch $\int_{\lambda}^{\lambda}$ ersetzen dürfen (vgl. I. 4., e) und Anm. ³¹⁾), und in diesem Integrationsbereich ist ($\lambda' \leq \lambda$), ist $\text{Min}(\lambda', \lambda) = \lambda'$, also gilt

$$(G(A)E(\lambda)f, g) = \int_{\lambda}^{\lambda} G(\lambda') d(E(\lambda')f, g).$$

Somit ist

$$\int H(\lambda) G(\lambda) \cdot d(E(\lambda)f, g) = \int H(\lambda) \cdot d\left[\int_{\lambda}^{\lambda} G(\lambda') \cdot d(E(\lambda')f, g)\right]$$

zu beweisen, und das ist gerade die Aussage von I. 4., i).

Nun können wir auch h) voll beweisen. Betrachten wir nämlich neben der Folge $F_n(x)$ und $F(x)$ auch die Folge $\overline{F_n(x)} F_n(x)$ und $\overline{F(x)} F(x)$, auf beide ist h) im bisher bewiesenen Umfange anwendbar, d. h. $F_n(A)$ strebt „schwach“ (in **B**) gegen $F(A)$, und $F_n(A)^* F_n(A)$ ebenso gegen $F(A)^* F(A)$ (hier sind die soeben bewiesenen b), e) heranzuziehen!). D. h. es ist für alle f, g $(F_n(A)f, g) \rightarrow (F(A)f, g)$ und $(F_n(A)^* F_n(A)f, g) \rightarrow (F(A)^* F(A)f, g)$, $(F_n(A)f, F_n(A)g) \rightarrow (F(A)f, F(A)g)$. Letzteres besagt für $f = g$ $|F_n(A)f| \rightarrow |F(A)f|$. Hieraus folgt aber bekanntlich $F_n(A)f \rightarrow F(A)f$ im Sinne der „starken“ Topologie (des Hilbertschen Raumes⁸⁷⁾). Die „Doppelkonvergenz“ ergibt sich schließlich, wenn wir die $\overline{F_n(x)}$ und $\overline{F(x)}$ betrachten.

Es bleibt noch übrig, f) zu beweisen. Nach g) können wir die darin vorkommenden $G(x)$, $H(x)$ (und mit ihnen $F(x) = G(H(x))$) auf gewissen Mengen abändern, ohne daß das Resultat beeinflußt wird, falls diese Menge die folgenden Eigenschaften haben: $H(x)$ werde nicht geändert, $G(x)$ auf einer Menge $\mathfrak{A}$. Damit $G(H(A))$ invariant bleibe, muß für jedes f die Menge der y , für die $\varrho(\psi_f(y))$ (man bilde es nach g) zu $H(A)$ zu $\mathfrak{A}$ gehört, eine Lebesguesche Nullmenge sein. Damit $F(A)$ invariant bleibe, muß für jedes f die Menge der y , für die $\varrho(\psi'_f(y))$ (man bilde es nach g) zu A zur Änderungsmenge von $F(x)$, d. h. $H(\psi'_f(y))$ zu $\mathfrak{A}$, gehört, eine Lebesguesche Nullmenge sein. Auf Grund der Bemerkung nach dem Beweise von g) genügt es dabei, jedesmal eine geeignete abzählbare Menge von f zu berücksichtigen. Daher ist Satz 3 anwendbar: wir dürfen $G(x)$ als zur dritten Baireschen Klasse gehörend annehmen, und wegen h) sogar als zur nullten gehörend — die ihrerseits als aus allen Polynomen bestehend angenommen werden kann. Also sei $G(x)$ ein Polynom.

Nun ist für $F(x) = a_0(H(x))^n + a_1(H(x))^{n-1} + \dots + a_n$ nach a)—d)

$$F(A) = a_0(H(A))^n + a_1(H(A))^{n-1} + \dots + a_n 1.$$

Und wenn wir $H(x) = x$ setzen, für $G(x) = a_0 x^n + a_1 x^{n-1} + \dots + a_n$

$$G(B) = a_0 B^n + a_1 B^{n-1} + \dots + a_n 1 \quad (\text{vgl. a}).$$

Daher gilt f) für alle Polynome $G(x)$.

3. Wir sind nun in der Lage, den folgenden Satz über die Konvergenz von Operatoren-Funktionen-Folgen zu beweisen:

Satz 4. Sei $F_1(x), F_2(x), \dots$ eine Folge Φ -meßbarer Funktionen und A ein Hermitescher Operator. Dann sind die folgenden Aussagen gleichbedeutend:

- a) Für jedes f von $\mathfrak{S}$ konvergieren die Punkte $F_1(A)f, F_2(A)f, \dots$ von $\mathfrak{S}$ gegen einen Limes, und zwar im Sinne der „starken“ Topologie von $\mathfrak{S}$ (vgl. a. a. O. Anm. ⁸⁵)).

β) α) gilt für die $F_1(A)f, F_2(A)f, \dots$, und für die $F_1(A)^*f, F_2(A)^*f, \dots$.
 γ) Alle $F_1(x), F_2(x), \dots$ sind in einer Menge $\mathfrak{M}$, die die in 2., g) angegebene Eigenschaft hat, gleichmäßig beschränkt. Wenn $m_1, m_2, \dots$ und $n_1, n_2, \dots$ zwei gegen Unendlich strebende Folgen ganzer Zahlen sind und f ein beliebiges Element von $\mathfrak{H}$, so hat die Folge $F_{m_1}(x) - F_{n_1}(x), F_{m_2}(x) - F_{n_2}(x), \dots$ eine Teilfolge $F_{m_{p_1}}(x) - F_{n_{p_1}}(x), F_{m_{p_2}}(x) - F_{n_{p_2}}(x), \dots$ ($p_1 < p_2 < \dots$), für welche die Menge $\mathfrak{M}$ aller x für die sie gegen 0 konvergiert, für dieses f die in 2., g) angegebene Eigenschaft hat.

Beweis: β) entsteht aus α) durch Hinzufügen derselben Forderung für $F_1(A)^*, F_2(A)^*, \dots$, d. h. für $\overline{F_1(x)}, \overline{F_2(x)}, \dots$. Da hierfür γ) ebenso lautet, genügt es α) und γ) zu vergleichen.

Betrachten wir zunächst den ersten Fall, wo γ) zutrifft. Wäre α) falsch, so gäbe es ein f , für welches die $F_n(A)f$ keinen starken Limes (in $\mathfrak{H}$) haben, d. h. (wegen der topologischen Vollständigkeit von $\mathfrak{H}$) das Cauchysche Konvergenzkriterium unerfüllt ist. Also: es gibt ein $\epsilon > 0$, so daß für beliebig große m geeignete $n \geq m$ existieren, für welche

$$|(F_m(A) - F_n(A))f| = |F_m(A)f - F_n(A)f| \geq \epsilon$$

ist. Daher lassen sich zwei gegen Unendlich strebende Folgen $m_1, m_2, \dots$ und $n_1, n_2, \dots$ finden, so daß stets $|(F_{m_\nu}(A) - F_{n_\nu}(A))f| \geq \epsilon$ ($\nu = 1, 2, \dots$) gilt. Insbesondere kann für die $p_1, p_2, \dots$ aus γ) nicht $(F_{m_{p_\nu}}(A) - F_{n_{p_\nu}}(A))f \rightarrow 0$ sein, obwohl dies wegen 2., h) der Fall sein sollte³⁸⁾.

Betrachten wir nun den zweiten Fall, wo α) zutrifft. Die erste Hälfte von γ) beweisen wir so: Da die Folge $F_n(A)f$ für jedes f konvergiert, sind die Operatoren $F_n(A)$ gleichmäßig beschränkt (vgl. A, Seite 382, der dort für stark konvergente $F_n(A)$ skizzierte, Banach-Sakssche, Beweis gilt ungeändert auch hier), d. h. es gibt eine Zahl C , so daß für alle f und n $|F_n(A)f| \leq C \cdot |f|$ gilt. Wenn wir also beweisen können, daß allgemein dies gilt: wenn für alle f $|G(A)f| \leq C \cdot |f|$ ist, so gilt $|G(x)| \leq C$ in einer Menge $\mathfrak{M}$ mit der in 2., g) angegebenen Eigenschaft — dann sind wir am Ziele (da die Durchschnittsmenge der zu $F_1(x), F_2(x), \dots$ derart gebildeten Mengen $\mathfrak{M}$ offenbar wieder die Eigenschaft aus 2., g) hat).

Für reelles $G(x)$, also Hermitesches $G(A)$, folgt aus $|G(A)f| \leq C \cdot |f|$ (für alle f) nach bekannten Sätzen von Hilbert, daß das Spektrum von $G(A)$ ganz im Intervalle $-C \leq x \leq C$ liegt. Wenn also eine Funktion $H(x)$ in diesem Intervalle überall verschwindet, so ist nach II., 2., g) und Anm. ³⁴⁾ $H(G(A)) = 0$, und daraus folgt durch Anwenden der Schlußbemerkung von II. auf $H(G(x))$, daß $H(G(x)) = 0$ in einer Menge $\mathfrak{M}$ vom erwähnten Charakter gilt. Setzen wir nun

$$H(x) = \begin{cases} 0 & \text{für } -C \leq x \leq C, \\ 1 & \text{sonst,} \end{cases}$$

so ist damit unsere Behauptung für reelle $G(x)$ bewiesen. Bei komplexem $G(x)$ haben wir $|G(A)f| \leq C \cdot |f|$, für $H(x) = \overline{G(x)}$ $|H(A)f| = |G(A)^*f| \leq C \cdot |f|$ (vgl. E., Anhang II., Seite 112), also für $K(x) = \Re G(x) = \frac{1}{2}(G(x) + H(x))$ $|K(A)f| \leq C \cdot |f|$. Somit gilt $|\Re G(x)| \leq C$ in einer $\mathfrak{M}$ -Menge. Betrachten wir statt $G(x)$ $\theta G(x)$ (θ konstant, $|\theta| = 1$), so ergibt sich ebenso: $|\Re(\theta G(x))| \leq C$ in einer $\mathfrak{M}$ -Menge. Wählen wir eine auf $|\theta| = 1$ überall dichte θ -Folge, so gilt obiges sogar für alle θ dieser Folge gleichzeitig im Durchschnitt einer $\mathfrak{M}$ -Mengenfolge, also in einer $\mathfrak{M}$ -Menge. Wo aber dies gilt, muß $|\Re(\theta G(x))| \leq C$ für alle θ mit $|\theta| = 1$ gelten — und dann ist $|G(x)| \leq C$.

Es bleibt übrig, die zweite Hälfte von γ) zu beweisen. Wäre sie falsch, so gäbe es zwei Folgen $m_1, m_2, \dots$ und $n_1, n_2, \dots$ sowie ein f , so daß für keine Teilfolge der Folge $F_{m_1}(x) - F_{n_1}(x), F_{m_2}(x) - F_{n_2}(x), \dots$ die Konvergenz gegen 0 in einer solchen Menge stattfindet, der alle $\varrho(\psi_f(y))$ angehören, bis auf eine Lebesguesche Nullmenge der y .

Nehmen wir an, daß für jedes $\varepsilon > 0, \eta > 0$ ein $\nu_0 = \nu_0(\varepsilon, \eta)$ existiert, so daß für jedes feste $\nu \geq \nu_0$ die Menge $\mathfrak{M}_\nu(\varepsilon)$ aller x , für die $|F_{m_\nu}(x) - F_{n_\nu}(x)| \geq \varepsilon$ ist, so beschaffen ist, daß alle $\varrho(\psi_f(y))$ zu ihr gehören, bis auf eine y -Menge von Lebesgueschem Maße $\leq \eta$. Dann sei $p_1 \geq \nu_0(\frac{1}{2}, \frac{1}{2})$, $p_2 \geq \nu_0(\frac{1}{4}, \frac{1}{4})$ und $> p_1, p_3 \geq \nu_0(\frac{1}{8}, \frac{1}{8})$ und $> p_2, \dots$. Das Zutreffen von $|F_{m_{p_\nu}}(x) - F_{n_{p_\nu}}(x)| < \frac{1}{2^\nu}$ bei allen $\nu \geq \mu$ erfolgt für alle x der Durchschnittsmenge

$$\mathfrak{M}_{p_\mu} \left(\frac{1}{2^\mu} \right) \cdot \mathfrak{M}_{p_{\mu+1}} \left(\frac{1}{2^{\mu+1}} \right) \cdot \dots,$$

welcher alle $\varrho(\psi_f(y))$ angehören, bis auf eine y -Menge von Lebesgueschem Maße $\leq \frac{1}{2^\mu} + \frac{1}{2^{\mu+1}} + \dots = \frac{1}{2^{\mu-1}}$. Für diese x konvergiert aber die Folge $F_{m_{p_\nu}}(x) - F_{n_{p_\nu}}(x)$ offenbar gegen 0; nennen wir die Konvergenzmenge der x $\mathfrak{M}$, so sehen wir: alle $\varrho(\psi_f(y))$ gehören zu $\mathfrak{M}$, bis auf eine y -Menge von Lebesgueschem Maße $\leq \frac{1}{2^{\mu-1}}$. Dies gilt für alle μ , also hat diese y -Menge das Lebesguesche Maß 0, entgegen dem früher Festgestellten. Daher existieren zwei $\varepsilon > 0, \eta > 0$, so daß für unendlich viele ν die Menge $\mathfrak{M}_\nu(\varepsilon)$ die $\varrho(\psi_f(y))$ nur bis auf eine y -Menge von Lebesgueschem Maße $\geq \eta$ enthält.

Für diese ν gilt aber:

$$\begin{aligned} & |(F_{m_\nu}(A) - F_{n_\nu}(A))f|^2 \\ &= ((F_{m_\nu}(A) - F_{n_\nu}(A))^2 f, f) \\ &= \int_{-\infty}^{\infty} (F_{m_\nu}(\xi) - F_{n_\nu}(\xi))^2 d|E(\xi)f|^2 \end{aligned}$$

$$\begin{aligned}
&= \int_0^\alpha (F_{m_\nu}(\varrho(x)) - F_{n_\nu}(\varrho(x)))^2 d|E(\varrho(x))f|^2 \\
&= \int_0^{|f|^2} \left(F_{m_\nu}(\varrho(\psi_f(y))) - F_{n_\nu}(\varrho(\psi_f(y))) \right)^2 dy \\
&\quad \varrho(\psi_f(y)) \text{ in } \mathfrak{M}_\nu(\varepsilon) \\
&\geq \int_{\varrho(\psi_f(y)) \text{ in } \mathfrak{M}_\nu(\varepsilon)} \left(F_{m_\nu}(\varrho(\psi_f(y))) - F_{n_\nu}(\varrho(\psi_f(y))) \right)^2 dy \\
&\geq \int \varepsilon^2 \cdot dy \geq \varepsilon^2 \eta.
\end{aligned}$$

Daher ist

$$\limsup_{\nu \rightarrow \infty} |F_{m_\nu}(A)f - F_{n_\nu}(A)f| \geq \sqrt{\varepsilon^2 \eta} > 0,$$

was mit der Konvergenz der $F_n(A)f$, d. h. α) unvereinbar ist. —

Zu diesem Satze bemerken wir noch: α) bzw. β) besagen, daß für jedes f $F_n(A)f$ einen starken Limes (in $\mathfrak{H}$) hat, bzw. $F_n(A)f$ und $F_n(A)^*f$ je einen solchen haben. In die Aussage von γ) geht dieses f auch ein. Satz 4 sagte nun aus, daß die Gültigkeit von α) für alle f der von β) für alle f und der von γ) für alle f gleichwertig ist — der Beweis zeigt aber, daß α), β) auch dann Folgen von γ) sind, wenn wir sie alle nur für ein bestimmtes f nehmen. Die gleichmäßige Beschränktheit der $F_n(x)$ in einer Menge $\mathfrak{M}$ mit der Eigenschaft aus 2., g) für alle f sei hier durchweg vorausgesetzt.) Da nun die $F_n(A)$, $F_n(A)^*$ gleichmäßig beschränkt sind (weil es die $F_n(x)$ im wesentlichen sind), ist die f -Menge in $\mathfrak{H}$, für die α) bzw. β) gilt, eine abgeschlossene lineare Mannigfaltigkeit — daher gelten sie überall, wenn sie etwa für ein vollständiges normiertes Orthogonalsystem gelten. Wegen der Gleichwertigkeit gilt dies auch für γ), d. h. γ) wird nicht abgeschwächt, wenn seine zweite Hälfte anstatt für alle f nur für die (abzählbar vielen) Elemente irgendeines vollständigen Orthogonalsystems gilt.

Wir beweisen nun weiter:

Satz 5. Seien die $F_n(x)$ und A wie in Satz 4. Wenn die dort erwähnte Konvergenz stattfindet, so gibt es einen Operator B aus $\mathbf{B}$, derart, daß die $F_n(A)f$ gegen Bf und die $F_n(A)^*f$ gegen B^*f konvergieren — d. h. B ist starker Doppellimes der $F_n(A)$ (vgl. a. a. O. Anm.⁸⁵). Dieses B ist wieder von der Form $F(A)$, wobei $F(x)$ auch beschränkt und Φ -meßbar ist.

Damit übrigens $F(A)$ starker (bzw. starker Doppel-) Limes der $F_n(A)$ sei, ist folgendes notwendig und hinreichend: Alle $F_n(x)$ sind in einer Menge $\mathfrak{M}$, die die in 2., g) angegebene Eigenschaft hat, gleichmäßig beschränkt. Wenn $n_1, n_2, \dots$ eine gegen Unendlich strebende Folge ist, und f ein beliebiges Element von $\mathfrak{H}$, so hat die Folge $F_{n_1}(x), F_{n_2}(x), \dots$ eine Teilfolge $F_{n_k}(x)$,

$F_{n_p}(x), \dots (p_1 < p_2 < \dots)$, für welche die Menge $\mathfrak{M}$ aller x , für die sie gegen $F(x)$ konvergiert, für dieses f die in 2., g) angegebene Eigenschaft hat.

Beweis: Die zweite Hälfte folgt sofort aus Satz 4, wenn wir die dortige Folge $F_1(x), F_2(x), \dots$ durch $F_1(x), F(x), F_2(x), F(x), \dots$ ersetzen. Die erste Hälfte folgt aus der zweiten, wenn wir eine Funktion $F(x)$ mit den in der zweiten Hälfte genannten Eigenschaften aufweisen können. Dabei ist aber zu bemerken, daß es wieder ausreicht, sich auf ein vollständiges normiertes Orthogonalsystem, d. h. eine abzählbare f -Menge $f_1, f_2, \dots$ zu beschränken — denn dies ist für die zweite Hälfte unseres Satzes ebenso zulässig, wie für Satz 4, γ), aus dem sie gewonnen wurde. Dabei ist unser Ausgangspunkt die Konvergenz der $F_n(A)$, also die in Satz 4, γ) genannten Eigenschaften der $F_n(x)$.

Sei zunächst f beliebig, aber fest, ferner $\varepsilon > 0, \eta > 0$. Es muß ein $\nu_0 = \nu_0(\varepsilon, \eta)$ geben, so daß für jedes gegebene System $m, n \geq \nu_0$ die Menge $\mathfrak{M}_{m,n}(\varepsilon)$ aller x mit $|F_m(x) - F_n(x)| \geq \varepsilon$ alle $\varrho(\psi_f(y))$ enthält, mit Ausnahme einer y -Menge von Lebesgueschem Maße $\leq \eta$. Im entgegengesetzten Falle gäbe es nämlich zwei gegen Unendlich strebende Folgen $m_1, m_2, \dots$ und $n_1, n_2, \dots$, so daß für jedes $\mathfrak{M}_{m_\nu, n_\nu}(\varepsilon)$ ($\nu = 1, 2, \dots$) die y mit nicht dazu gehörendem $\varrho(\psi_f(y))$ ein Lebesguesches Maß $\geq \eta$ haben. Aus der m_ν, n_ν -Folge wählen wir die in Satz 4, γ) beschriebene m_{p_ν}, n_{p_ν} -Teilfolge ($p_1 < p_2 < \dots$) aus. Da nun alle in Frage kommenden y in einer Menge endlichen Maßes (nämlich $0 \leq y \leq |f|^2$, der Wertevorrat von $\Phi_f(x) = |E(\varrho(x))f|^2$) enthalten sind, ist der Satz von Arzelà-Young³⁹⁾ anwendbar: die Menge der y , für die $\varrho(\psi_f(y))$ für unendlich viele ν dem $\mathfrak{M}_{m_{p_\nu}, n_{p_\nu}}(\varepsilon)$ nicht angehört, hat auch ein Maß $\geq \eta$. Für ein solches $x = \varrho(\psi_f(y))$ ist also unendlich oft $|F_{m_{p_\nu}}(x) - F_{n_{p_\nu}}(x)| \geq \varepsilon$, also gewiß nicht $F_{m_{p_\nu}}(x) - F_{n_{p_\nu}}(x) \rightarrow 0$. Nach Satz 4, γ) sollten aber diejenigen y , deren $x = \varrho(\psi_f(y))$ eine nicht gegen 0 konvergente Folge $F_{m_{p_\nu}}(x) - F_{n_{p_\nu}}(x)$ erzeugen, eine Menge vom Lebesgueschen Maße 0 bilden. Dies ist ein Widerspruch.

Nun sei $n_1, n_2, \dots$ eine gegen Unendlich strebende Folge. Wir wählen p_1 mit $n_{p_1} \geq \nu_0 \left(\frac{1}{2}, \frac{1}{2}\right)$, $p_2 > p_1$ mit $n_{p_2} \geq \nu_0 \left(\frac{1}{2}, \frac{1}{2}\right)$, $\nu_0 \left(\frac{1}{4}, \frac{1}{4}\right)$, $p_3 > p_2$ mit $n_{p_3} \geq \nu_0 \left(\frac{1}{4}, \frac{1}{4}\right)$, $\nu_0 \left(\frac{1}{8}, \frac{1}{8}\right), \dots$. Daher gilt $|F_{n_{p_\nu}}(x) - F_{n_{p_{\nu+1}}}(x)| \leq \frac{1}{2^\nu}$ für alle $\varrho(\psi_f(y))$, mit Ausnahme einer y -Menge vom Lebesgueschen Maße $\leq \frac{1}{2^\nu}$. Es gilt also für alle $\nu \geq \mu$ (μ gegeben) zugleich, mit Ausnahme einer y -Menge vom Lebesgueschen Maße $\leq \frac{1}{2^\mu} + \frac{1}{2^{\mu+1}} + \dots = \frac{1}{2^{\mu-1}}$. Für diese y , bzw. ihre $x = \varrho(\psi_f(y))$, ist aber dann für alle $\nu \geq \mu, \nu < \varrho$

$$|F_{n_{p_\nu}}(x) - F_{n_{p_\rho}}(x)| \leq \frac{1}{2^\nu} + \frac{1}{2^{\nu+1}} + \cdots + \frac{1}{2^{\rho-1}} < \frac{1}{2^{\rho-1}}.$$

Also erfüllen die $F_{n_{p_\nu}}(x)$ das Cauchysche Konvergenzkriterium, d. h. die Folge $F_{n_{p_\nu}}(x)$ ($\nu = 1, 2, \dots$) konvergiert. Das Maß der Menge der y , für deren $x = \varrho(\psi_f(y))$ sie divergiert, ist somit $\leq \frac{1}{2^{\mu-1}}$ — und da dies für alle μ gilt, ist sie $= 0$.

Nunmehr wollen wir f variieren. Wir nehmen eine gegen Unendlich strebende Folge $n_1, n_2, \dots$ und gehen wie folgt vor: Aus $F_{n_1}(x), F_{n_2}(x), \dots$ wählen wir nach obigem eine Teilfolge $F_{r_1}(x), F_{r_2}(x), \dots$ aus, die für alle $\varrho(\psi_{f_1}(y))$ konvergiert, mit Ausnahme einer y -Menge vom Maße 0. Aus $F_{r_2}(x), F_{r_3}(x), \dots$ wählen wir eine Teilfolge $F_{s_1}(x), F_{s_2}(x), \dots$ aus, die für alle $\varrho(\psi_{f_2}(y))$ konvergiert, mit Ausnahme einer y -Menge vom Maße 0. Aus $F_{s_2}(x), F_{s_3}(x), \dots$ wählen wir eine Teilfolge $F_{t_1}(x), F_{t_2}(x), \dots$ aus, die für alle $\varrho(\psi_{f_3}(y))$ konvergiert, mit Ausnahme einer y -Menge vom Maße 0, Die Folge $F_{r_1}(x), F_{s_1}(x), F_{t_1}(x), \dots$ ist nun eine Teilfolge von $F_{n_1}(x), F_{n_2}(x), \dots$, und für jede der Folgen $F_{r_1}(x), F_{r_3}(x), \dots, F_{s_1}(x), F_{s_2}(x), \dots, F_{t_1}(x), F_{t_2}(x), \dots, \dots$ eine Teilfolge nach Wegnahme endlich vieler Glieder. Die Menge aller x , für die sie konvergiert, heiße $\mathfrak{M}$, dann enthält wegen der letzten Bemerkung $\mathfrak{M}$ alle $\varrho(\psi_{f_n}(y))$ mit Ausnahme einer y -Menge vom Maße 0, und zwar für jedes n . Statt $F_{r_1}(x), F_{s_1}(x), \dots$ schreiben wir lieber $F_{n_{p_1}}(x), F_{n_{p_2}}(x), \dots$ ($p_1 < p_2 < \dots$), die Limesfunktion nennen wir $F(x)$ — diese ist vorläufig nur in $\mathfrak{M}$ definiert, außerhalb $\mathfrak{M}$ möge sie etwa 0 sein. Da die $F_{n_{p_\nu}}(x)$ gleichmäßig beschränkt sind, ist es auch $F(x)$; da sie Φ -meßbar sind, sind es auch $\mathfrak{M}$ und $F(x)$ (man sieht dies genau so ein, wie für die gewöhnliche Meßbarkeit).

Auf Grund unserer Konstruktion scheint $F(x)$ noch von der Folge $n_1, n_2, \dots$ abzuhängen. Seien $m_1, m_2, \dots$ und $n_1, n_2, \dots$ zwei solche Folgen, und $m_{p_1}, m_{p_2}, \dots$ bzw. $n_{q_1}, n_{q_2}, \dots$ die nach obigem hergestellten Teilfolgen, derart, daß $F_{m_{p_\nu}}(x)$ für $\nu \rightarrow \infty$ gegen $F'(x)$ konvergiert und $F_{n_{q_\nu}}(x)$ für $\nu \rightarrow \infty$ gegen $F''(x)$ — beides für alle $x = \varrho(\psi_{f_l}(y))$ mit Ausnahme einer y -Menge vom Maße 0, und zwar für jedes l . Nach Satz 4, γ), gibt es nun, wenn wir l für den Augenblick festhalten, eine (von l abhängige) Folge $u_1 < u_2 < \dots$, so daß $F_{m_{p_{u_\nu}}}(x) - F_{n_{q_{u_\nu}}}(x) \rightarrow 0$ für alle $\varrho(\psi_{f_l}(y))$ stattfindet, mit Ausnahme einer y -Menge vom Maße 0. Daher ist $F'(x) = F''(x)$, für alle $\psi_f(y)$, mit Ausnahme einer y -Menge vom Maße 0, und dies gilt für alle l . D. h.: $F(x)$ ist im wesentlichen von der Folge $n_1, n_2, \dots$ unabhängig, wenn wir es für eine solche Folge $n_1, n_2, \dots$ wählen und festhalten, wird es doch auch für alle anderen Folgen, die in der zweiten

Hälfte unseres Satzes auseinandergesetzten Eigenschaften haben. Damit haben wir alles nachgeholt, was notwendig war, um auch die erste Hälfte unseres Satzes zu beweisen. —

Ehe wir diesen Teil abschließen, sei noch erwähnt, daß weder Satz 4 noch Satz 5 für die „schwache“ Topologie von $\mathbf{B}$ (vgl.-a. a. O. Anm. ³⁵)) zutrifft, die einfachen diesbezüglichen Gegenbeispiele übergehen wir. Ferner bemerken wir, daß ein notwendiges und hinreichendes Kriterium für das Bestehen von $F(A) = G(A)$ leicht aus der zweiten Hälfte von Satz 5 ableitbar ist: es genügt, alle $F_n(x)$ gleich $G(x)$ zu setzen. Dann finden wir: $F(x) = G(x)$ muß für jedes f für alle $\varphi(\Psi_f(y))$ zutreffen, mit Ausnahme einer y -Menge vom Lebesgueschen Maße 0, d. h. die hinreichende Bedingung in 2, g) ist auch notwendig. Ferner genügt es, ebenso wie in Satz 5, dies anstatt für alle f , nur für ein vollständiges normiertes Orthogonalsystem zu verlangen. Schließlich sei noch Folgendes hervorgehoben: In der zweiten Hälfte von Satz 4, γ) sowie von Satz 5 wird die dort auftretende Teilfolge $F_{m_\nu}(x) - F_{n_\nu}(x)$ bzw. $F_{n_\nu}(x) - F(x)$ ($\nu = 1, 2, \dots$) in Abhängigkeit von f gewählt, und ihr $\mathfrak{M}$ hat die Eigenschaft aus 2, g) nur für dieses f . Wir können aber an beiden Stellen ebensogut verlangen, daß sie von f unabhängig wählbar sei — ihr festes $\mathfrak{M}$ hat dann die Eigenschaft aus 2, g) allgemein. Sei nämlich $f_1, f_2, \dots$ ein vollständiges normiertes Orthogonalsystem, ein Hilbertsches Diagonalverfahren, genau wie beim Beweise von Satz 5, liefert dann eine Teilfolge, die (d. h. deren $\mathfrak{M}$) das Gewünschte für alle f aus $f_1, f_2, \dots$ tut. Daß aber ein solches $\mathfrak{M}$ das Gewünschte (d. i. das Erfülltsein der Eigenschaft aus 2, g) für alle f überhaupt leistet, folgt z. B. aus dem obigen Kriterium für $F(A) = G(A)$, wo es auch gleichgültig war, ob es für das dortige (offenbar allgemeine) $\mathfrak{M}$ für alle f oder nur für die $f_1, f_2, \dots$ verlangt wird.

III. Operatorenfunktionen und -ringe.

1. Die Sätze 4, 5 ermöglichen einen einfachen Beweis der in Einleitung, 3, aufgestellten Behauptung.

Satz 6. *Sei A ein beschränkter Hermitescher Operator, $\mathbf{R}(A)$ sein Ring. Dann ist $\mathbf{R}(A)$ die Menge aller $F(A)$, wobei $F(x)$ alle Φ -meßbaren Funktionen mit $F(0) = 0$ durchläuft. (Ob die $F(x)$ als beschränkt, oder, wie in Anm. ²⁸) bloß als in jedem endlichen Intervall beschränkt, postuliert werden, ist gleichgültig.)*

Beweis: Zunächst gehören alle diese $F(A)$ zu $\mathbf{R}(A)$. Denn es genügt, wie wir mehrmals schlossen, die $F(x)$ der dritten Baireschen Klasse zu berücksichtigen, und zwar natürlich die mit $F(0) = 0$. Sodann genügt es, wegen der Geschlossenheitseigenschaften der Ringe (und II., 2., h), oder Satz 5.), sogar die der nullten Baireschen Klasse zu betrachten, als welche

hier die Polynome $p(x)$ mit $p(0) = 0$ fungieren dürfen. Daß aber ein solches $p(A)$ zum Ringe von A gehört, wissen wir längst, z. B. aus A , Seite 404.

Um die Umkehrung zu beweisen, erinnern wir daran, daß nach der Bemerkung a. a. O. jeder Operator von $\mathbf{R}(A)$ Limes einer stark doppelkonvergenten Folge $p_1(A), p_2(A), \dots$ (alle $p_n(x)$ Polynome, $p_n(0) = 0$) ist. Aus Satz 5 folgt daher, daß er die Form $F(A)$ hat, und aus der dort gegebenen Charakterisierung von $F(x)$, daß dieses beschränkt und $\mathcal{O}$ -meßbar ist, sowie die Möglichkeit, es mit $F(0) = 0$ zu wählen.

Damit ist alles bewiesen.

Zwei Bemerkungen sind naheliegend: Erstens, daß es genügt, sich auf die $F(x)$ der dritten Baireschen Klasse zu beschränken. Zweitens, daß die Beschränkung $F(0) = 0$ fallen gelassen werden kann, falls eine Abänderung von $F(x)$ im Punkte $x = 0$ $F(A)$ nicht berührt. D. h., wenn für jedes f $\varrho(\psi_f(y)) = 0$ nur für eine Lebesguesche Nullmenge eintritt, d. h. wenn 0 nicht der Wert dieser Funktion in einem Konstanzintervall ist. Also: $\varrho^{-1}(0)$ nicht ein Konstanzintervallwert von $\psi_f(y)$, oder: $\varrho^{-1}(0)$ keine Unstetigkeitsstelle von $\mathcal{O}_f(x) = |E(\varrho(x))f|^2$, oder: 0 keine Unstetigkeitsstelle von $|E(\xi)f|^2$. D. h.: wenn 0 nicht zum Punktspektrum $\mathfrak{P}$ von A gehört, oder: wenn $Af = 0$ $f = 0$ zur Folge hat.

Die Hermiteschen Operatoren des Ringes $\mathbf{R}(A)$ erhalten wir, wenn wir uns auf die reellwertigen $F(x)$ beschränken. Denn daß für reelles $F(x)$ $F(A)$ Hermitesch ist, ist klar; ist umgekehrt $F(A)$ Hermitesch, so ist es gleich $F(A)^*$, also gleich $\frac{F(A) + F(A)^*}{2} = G(A)$, wo $G(x)$ die reelle Funktion $\frac{F(x) + \overline{F(x)}}{2}$ ist.

Da eine Abelsche Menge $\mathbf{M}$ von Operatoren aus $\mathbf{B}$ (d. h. eine Menge, für deren irgend zwei Elemente A, B A mit B und B^* bei der Multiplikation vertauschbar ist — daher sind alle ihre Elemente normal) stets Teil eines Abelschen Ringes ist (vgl. A , Seite 389), also, nach dem Hauptresultat von A (Seite 401, Satz 10), eines Ringes $\mathbf{R}(A)$, gibt es gewiß einen Hermiteschen Operator A , von dem alle Elemente von $\mathbf{M}$ Funktionen sind. Besteht $\mathbf{M}$ aus Hermiteschen Operatoren, so ist es Abelsch, wenn diese alle untereinander bei der Multiplikation vertauschbar sind, und dann sind die obengenannten Funktionen sogar reell wählbar.

Noch eines ist erwähnenswert: Sei A ein Hermitescher Operator, dann betrachten wir diejenigen Operatoren B (aus $\mathbf{B}$), die mit jedem Operator vertauschbar sind, der mit A vertauschbar ist. In A , Seite 404, wurde bewiesen, daß diese B alle normal sind, und zwar sind es gerade die Operatoren $a1 + B'$, wo a alle komplexen Zahlen und B' alle Elemente

von $\mathbf{R}(A)$ durchläuft. Nach Satz 6 sind es also alle $a1 + F'(A) = F(A)$ ($F(x) = a + F'(x)$), wo a alle komplexen Zahlen und $F'(x)$ die Funktionen mit $F'(0) = 0$ durchläuft — also $F(x)$ alle Funktionen. Wenn wir B als Hermitesch voraussetzen, so genügt es wiederum, die reellen $F(x)$ zu betrachten.

2. Sei $\mathbf{R}$ ein Abelscher Ring, dieser kann dann stets als $\mathbf{R}(A)$ geschrieben werden (vgl. A, Seite 401, Satz 10), A Hermitesch — aber A ist hierdurch, wie man leicht erkennt, noch lange nicht festgelegt. Sei also $\mathbf{R} = \mathbf{R}(A) = \mathbf{R}(B)$, welcher Zusammenhang muß dann zwischen A und B (beide seien Hermitesch) bestehen?

Da jedes zum Ringe des andern gehört, ist jedes Funktion, und zwar reelle Funktion, des anderen: $B = F(A)$, $A = G(B)$, und dabei $F(0) = 0$, $G(0) = 0$. Dies ist übrigens offenbar auch hinreichend, damit A zum Ringe von B und B zum Ringe von A gehöre, d. h. damit $\mathbf{R}(A) = \mathbf{R}(B)$ sei. Wir betrachten nun die (natürlich in jedem endlichen Intervalle beschränkte und $\mathcal{O}$ -meßbare) Funktion $F(x)$ (mit $F(0) = 0$) als gegeben, und fragen: wie muß sie beschaffen sein, damit ein (ebenfalls in jedem endlichen Intervalle beschränktes und $\mathcal{O}$ -meßbares) $G(x)$ (mit $G(0) = 0$) existiert, so daß für den durch $B = F(A)$ definierten Operator B $A = G(B)$ gilt?

Diese Forderung besagt $G(F(A)) = A$, also, nach der Schlußbemerkung von II., 3., $G(F(x)) = x$ in einer x -Menge $\mathfrak{M}$ mit der in II., 2., g) angegebenen Eigenschaft. Wie man sieht, kommt dies darauf heraus, daß das $\mathcal{O}$ -meßbare $F(x)$ eine $\mathcal{O}$ -meßbare Inverse $G(x)$ haben soll — mit Ausnahme einer x -Menge, für die die Urbilder aller Abbildungen $\varphi(\mathcal{U}_f(y)) = x$ Nullmengen sind (vgl. II., 2., g)). (Die Beschränktheit von $G(x)$ ist unwesentlich: es kommt ja nur auf die x des Spektrums von A an, also in einem endlichen Intervalle, d. h. soweit der Wertverlauf von $G(u)$ überhaupt vorgeschrieben ist, liegt er zwischen endlichen Schranken.) Es wäre nahelegend, sich von der Annahme frei zu machen zu suchen, also etwa mit solchen $F(x)$ zu operieren, die jeden (reellen) Wert einmal und nur einmal annehmen. Indessen müßte dann auch noch die weitere Frage klargestellt werden, wann die Inverse einer $\mathcal{O}$ -meßbaren Funktion (die eine Inverse besitzt) auch $\mathcal{O}$ -meßbar ist? Diese Frage scheint aber ziemlich kompliziert zu sein⁴⁰⁾.

IV. Mehrvariablen-Funktionen und normale Operatoren.

1. Zunächst führen wir den Begriff der $\mathcal{O}$ -meßbarkeit für derartige (komplexwertige) Funktionen $F(x_1, x_2, \dots)$ ein, die von einer endlichen oder abzählbar unendlichen Zahl komplexer Variablen $x_1, x_2, \dots$ abhängen (auch bei einer endlichen Zahl von Variablen, wo es $x_1, \dots, x_n$

heißen sollte, schreiben wir der Einheitlichkeit halber $x_1, x_2, \dots$). Ein solches $F(x_1, x_2, \dots)$ nennen wir Φ -meßbar, wenn es die folgende Eigenschaft hat: Wenn $f_1(u), f_2(u), \dots$ ein beliebiges System für alle reellen u definierter und nach rechts halbstetiger Funktionen ist⁴¹⁾, so ist $F(f_1(u), f_2(u), \dots)$ als komplexwertige Funktion des reellen u aufgefaßt!) meßbar. Offenbar sind alle Funktionen, die nur von endlich vielen Variablen abhängen, und in diesen stetig sind, Φ -meßbar; und wenn $F_1, F_2, \dots$ eine Folge Φ -meßbarer Funktionen gilt, die überall gegen ein F konvergieren, so ist auch F Φ -meßbar. Statt dieses letzteren beweisen wir gleich mehr: Seien $F_1, F_2, \dots$ Φ -meßbar, und F so definiert:

$$F(x_1, x_2, \dots) = \begin{cases} \lim_{n \rightarrow \infty} F_n(x_1, x_2, \dots), & \text{wenn dieser Limes existiert} \\ 0 & \text{sonst.} \end{cases}$$

Da nämlich Φ -Meßbarkeit von F_n, F die gewöhnliche Meßbarkeit von $F_n(f_1(u), f_2(u), \dots), F(f_1(u), f_2(u), \dots)$ (bei jeder Wahl der $f_1(u), f_2(u), \dots$) bedeutet, genügt es, den obigen Satz für Funktionen einer reellen Variablen u und gewöhnliche Meßbarkeit zu beweisen — dort ist er aber bekanntlich richtig. Somit sind alle Funktionen aller Baireschen Klassen wiederum Φ -meßbar⁴²⁾. Aus den analogen Eigenschaften der Meßbarkeit für Funktionen einer (reellen) Variablen folgt wieder unmittelbar, daß mit F, G auch $F \pm G, aF, FG$ Φ -meßbar sind. Etwas tiefer liegt der folgende Satz:

SATZ 7. Wenn $F(x_1, x_2, \dots)$ Φ -meßbar ist und $G_1(y_1, y_2, \dots), G_2(y_1, y_2, \dots), \dots$ meßbar bzw. Φ -meßbar (alle Funktionen komplexwertig und mit komplexen Argumenten), so ist auch $F(G_1(y_1, y_2, \dots), G_2(y_1, y_2, \dots), \dots)$ meßbar bzw. Φ -meßbar. (Man beachte, daß die erste, die gewöhnliche Meßbarkeit betreffende, Behauptung nur bei endlicher Anzahl der Variablen einen Sinn hat — da die gewöhnliche Meßbarkeit, im Gegensatz zur Φ -Meßbarkeit, nur dann definiert ist.)

Beweis: Die zweite Behauptung (Φ -Meßbarkeit) läuft darauf heraus, daß für alle nach rechts halbstetigen $f_1(u), f_2(u), \dots$ die Funktion

$$F(G_1(f_1(u), f_2(u), \dots), G_2(f_1(u), f_2(u), \dots), \dots)$$

meßbar ist — da aber die $G_1(f_1(u), f_2(u), \dots), G_2(f_1(u), f_2(u), \dots), \dots$ meßbar sind, ist das die erste Behauptung, für eine reelle Variable. Die erste Behauptung für endlich viele komplexe Variable $y_1, \dots, y_n$ (dies ist ihre allgemeine Form) kommt auf dasselbe heraus, wie für doppelt so viele reelle $\Re y_1, \Im y_1, \dots, \Re y_n, \Im y_n$ — somit ist lediglich die erste Behauptung für endlich viele reelle Variable zu beweisen. Diese reellen Variablen nennen wir wieder $y_1, y_2, \dots$. Indem wir ferner die Variablen $x_1, x_2, \dots$ von F durch $\Re x_1, \Im x_1, \Re x_2, \Im x_2, \dots$ ersetzen, und entsprechend die

$G_1, G_2, \dots$ durch $\Re(G_1), \Im(G_1), \Re(G_2), \Im(G_2), \dots$, erreichen wir, daß auch F reelle Argumente hat — auch diese nennen wir wieder $x_1, x_2, \dots$, und die einzusetzenden Funktionen wieder $G_1, G_2, \dots$. Schließlich können wir, durch denselben Kunstgriff, wie bei Beginn des Beweises von Satz 2, alle Variablen $x_1, x_2, \dots$ und $y_1, \dots, y_n$ auf endliche Intervalle beschränken. Anstatt des dortigen Intervalles $0 < \xi < 2$ wählen wir jetzt $0 < \xi < 1$, daß natürlich zu $0 \leq \xi < 1$ erweitert werden kann. Also: $F(x_1, x_2, \dots)$ ist in $0 \leq x_1 < 1, 0 \leq x_2 < 1, \dots$ definiert, die $G_1(y_1, \dots, y_n), G_2(y_1, \dots, y_n), \dots$ in $0 \leq y_1 < 1, \dots, 0 \leq y_n < 1$, und die Werte der $G_1, G_2, \dots$ sind (da sie in F substituiert werden), stets $\geq 0, < 1$.

Nehmen wir nun an, der Beweis wäre für $n = 1$ geführt. Der allgemeine Fall mit $y_1, \dots, y_n$ wird dann folgendermaßen auf denjenigen mit nur einem y zurückgeführt: Sei $y = \alpha(y_1, \dots, y_n), y_1 = \beta_1(y), \dots, y_n = \beta_n(y)$ eine ein-eindeutige Abbildung des Intervalles $0 \leq y < 1$ auf den Würfel $0 \leq y_1 < 1, \dots, 0 \leq y_n < 1$ (bei beiden darf je eine, lineare bzw. n -dimensionale Lebesguesche Nullmenge ausgelassen werden), die maßtreu ist, d. h.: y -Urbild und $\{y_1, \dots, y_n\}$ -Bild sind beide meßbar oder beide unmeßbar, und haben im ersteren Falle dasselbe Maß (im linearen bzw. im n -dimensionalen Sinne)⁴³. Dann ist mit $G_1(y_1, \dots, y_n), G_2(y_1, \dots, y_n), \dots$ auch $G_1(\alpha_1(y), \dots, \alpha_n(y)), G_2(\alpha_1(y), \dots, \alpha_n(y)), \dots$ meßbar, also, wenn der Satz für eine Variable feststeht, auch $F(G_1(\alpha_1(y)), \dots, \alpha_n(y)), G_2(\alpha_1(y), \dots, \alpha_n(y)), \dots)$, und dies bleibt meßbar, wenn darin y durch $\beta(y_1, \dots, y_n)$ ersetzt wird — d. h. $F(G_1(y_1, \dots, y_n), G_2(y_1, \dots, y_n), \dots)$ ist meßbar. Wir gehen daher zur Erledigung des Falles mit einer einzigen reellen Variablen y über.

Seien $f_1(u), f_2(u), \dots$ rechts halbstetig, $\varphi(y)$ monoton nichtfallend und halbstetig. Dann sind alle $f_1(\varphi(y)), f_2(\varphi(y)), \dots$ rechts halbstetig, also ist $F(f_1(\varphi(y)), f_2(\varphi(y)), \dots)$ meßbar. Da dies für alle genannten $\varphi(y)$ gilt, ist $F(f_1(u), f_2(u), \dots)$ $\mathcal{O}$ -meßbar, also nach Satz 2 für jedes meßbare $G(y)$ $F(f_1(G(y)), f_2(G(y)), \dots)$ meßbar. Wenn nun zu jedem System $G_1(y), G_2(y), \dots$ meßbarer Funktionen eine meßbare Funktion $G(y)$ und rechts halbstetige $f_1(u), f_2(u), \dots$ gefunden werden können, so daß $G_1(y) = f_1(G(y)), G_2(y) = f_2(G(y)), \dots$ ist, so sind wir demnach am Ziele. Diese Konstruktion führen wir nun so durch: Sei $x = \alpha(x_1, x_2, \dots), x_1 = \beta_1(x), x_2 = \beta_2(x), \dots$ eine ein-eindeutige Abbildung einer Teilmenge des Intervalles $0 \leq x < 1$ auf den (ganzen) Würfel $0 \leq x_1 < 1, 0 \leq x_2 < 1, \dots$. Dabei seien alle $\beta_1(x), \beta_2(x), \dots$ nach rechts halbstetig, und sie mögen die folgende Eigenschaft besitzen: wenn x den in einem Intervalle $\frac{p}{2^q} \leq x < \frac{p+1}{2^q}$ ($q = 0, 1, \dots, p = 0, 1, \dots, 2^q - 1$) gelegenen Teil der Bildmenge durchläuft, so durchläuft der Punkt $x_1 = \beta_1(x), x_2 = \beta_2(x), \dots$

einen Würfel (d. i. eine Menge $a_1 \leq x_1 < b_1, a_2 \leq x_2 < b_2, \dots$)⁴⁴). Betrachten wir nun $G(y) = \alpha(G_1(y), G_2(y), \dots)$. $\frac{p}{2^q} \leq G(y) < \frac{p+1}{2^q}$ kommt $a_1 \leq G_1(y) < b_1, a_2 \leq G_2(y) < b_2, \dots$ gleich, findet also in einer y -Menge statt, die Durchschnitt abzählbar vieler meßbarer Mengen ist (die $G_1(y), G_2(y), \dots$ sind ja meßbar), folglich selbst meßbar ist. Nun ist die y -Menge von $G(y) < a$ (a beliebig) als Summe abzählbar vieler solcher Mengen herstellbar, also auch meßbar — somit ist $G(y)$ meßbar. Wenn wir also $f_1 \equiv \beta_1, f_2 \equiv \beta_2, \dots$ setzen, so sind wir am Ziele.

Aus Satz 7 folgt sofort: ist die Zahl der $x_1, x_2, \dots$ endlich, so ist ein Φ -meßbares $F(x_1, x_2, \dots)$ meßbar — es genügt, die $y_1, y_2, \dots$ mit den $x_1, x_2, \dots$ zusammenfallen zu lassen, und $G_1(y_1, y_2, \dots) = y_1, G_2(y_1, y_2, \dots) = y_2, \dots$ zu setzen. Der Vergleich des neuen Φ -Meßbarkeits-Begriffes mit dem alten (I., 3.) ist auch bei einer einzigen Variablen x unmöglich, da diese jetzt komplex ist, und früher reell war. Hängt aber $F(x)$ etwa nur von $\Re(x)$ ab, so sind beide gleichwertig: nach Definition beider folgt die alte Φ -Meßbarkeit aus der neuen, nach Satz 2 und der Definition der neuen Φ -Meßbarkeit gilt aber auch die Umkehrung.

Wenn $F(x_1, x_2, \dots)$ im neuen Sinne Φ -meßbar ist, und $G_1(x), G_2(x), \dots$ im alten (x reell!), so sind es, nach dem soeben Gesagten, $G_1(\Re x), G_2(\Re x), \dots$ (x komplex!) auch im neuen Sinne. Nach Satz 7 ist es daher $F(G_1(\Re x), G_2(\Re x), \dots)$ auch, und für reelles x ist es $F(G_1(x), G_2(x), \dots)$ wieder im alten Sinne.

Auf Grund des bisher Gesagten können wir nunmehr die Unterscheidung zwischen beiden Definitionen der Φ -Meßbarkeit ohne Verwechslungsgefahr fallen lassen.

2. Sei nun $A_1, A_2, \dots$ eine endliche oder abzählbar unendliche Abelsche Menge linear-beschränkter Operatoren (vgl. A, Seite 389), d. h. es soll jedes A_m mit jedem A_n und A_n^* vertauschbar sein — folglich sind insbesondere alle A_n normal. $F(x_1, x_2, \dots)$ sei eine Φ -meßbare Funktion mit so vielen Variablen, als die Zahl der $A_1, A_2, \dots$ beträgt.

Nach einer der Bemerkungen in III., 1. existieren solche (beschränkte) Hermitesche Operatoren, von denen alle $A_1, A_2, \dots$ Funktionen sind. Sei etwa A ein solcher, also $A_1 = G_1(A), A_2 = G_2(A), \dots, G_1(x), G_2(x), \dots$ alle Φ -meßbar (x reell). Die Funktion $H(x) = F(G_1(x), G_2(x), \dots)$ ist somit auch Φ -meßbar, und wenn sie in einer Menge $\mathfrak{M}$ mit den in II., 2., g) angegebenen Eigenschaften (mit dem Operator A) beschränkt ist, so können wir $H(A)$ bilden — dies ist ein beschränkt-linearer, normaler Operator. Ihn wollen wir als $F(A_1, A_2, \dots)$ definieren, müssen aber zu diesem Zwecke zunächst zeigen, daß diese Definition vom hilfsweise eingeführten A nur scheinbar abhängt: d. h. daß sie für alle genannten A ausführbar ist, oder

für keines, und daß sie im ersteren Falle für alle dasselbe Resultat ergibt.

Der Ring $\mathbf{R}(A_1, A_2, \dots)$ ist Abelsch, also nach A, Seite 401, Satz 10, ein $\mathbf{R}(B)$ mit Hermiteschem B ; nach Satz 6 sind somit alle $A_1, A_2, \dots$ Funktionen von B — somit ist B als ein A im Sinne von vorhin verwendbar. Wir wollen nun zeigen, daß sich jedes der genannten A für die Zwecke unserer Definition genau so benimmt, wie dieses B , und zum selben Resultate führt — damit ist dann alles Erforderliche geschehen.

Es ist $G_1(A) = A_1$, $G_2(A) = A_2, \dots$ und $G'_1(B) = A_1$, $G'_2(B) = A_2, \dots$. $\mathbf{R}(A)$ enthält alle Funktionen von A , also auch $A_1, A_2, \dots$ also umfaßt es $\mathbf{R}(A_1, A_2, \dots) = \mathbf{R}(B)$, also enthält es B . Daher ist $B = K(A)$. Hieraus folgt $G_1(A) = G'_1(K(A))$, $G_2 = G'_2(K(A))$, $\dots$, und nach II., 2., f) sowie der Schlußbemerkung von II., 3. $G_1(x) = G'_1(K(x))$ in einer Menge $\mathfrak{M}$ mit den Eigenschaften II., 2., g) (für A), $G_2(x) = G'_2(K(x))$ in einer eben solchen Menge, $\dots$. Alle diese Gleichungen gelten gleichzeitig im Durchschnitt dieser Mengenfolge, d. h. wieder in einer Menge $\mathfrak{M}$ nach II., 2., g). Wir haben nun im Sinne unserer Definition bei A $H(x) = F(G_1(x), G_2(x), \dots)$ zu bilden, bei B $H'(x) = F(G'_1(x), G'_2(x), \dots)$. Nach obigem ist $H(x) = H'(K(x))$ in einer Menge $\mathfrak{M}$ nach II., 2., g). Aus II., 2., f) folgt sofort, daß $H'(B) = H'(K(A)) = H(A)$ ist, d. h. daß beide Definitionswege dasselbe ergeben, falls $H(x)$ für A und $H'(x)$ für B die gewünschten Beschränktheitseigenschaften haben. Es bleibt übrig, die Gleichwertigkeit dieser zu beweisen; wir beweisen aber lieber mehr: die Gültigkeit von $|H(x)| \leq C$ in einem $\mathfrak{M}$ nach II., 2., g) für A ist derjenigen von $|H'(x)| \leq C$ in einem $\mathfrak{M}$ nach II., 2., g) für B gleichwertig.

Ersetzen wir $F(x_1, x_2, \dots)$ durch

$$F_D(x_1, x_2, \dots) = \begin{cases} F(x_1, x_2, \dots) & \text{für } |F(x_1, x_2, \dots)| \leq D, \\ D & \text{sonst,} \end{cases}$$

dann gehen $H(x)$, $H'(x)$ in

$$H_D(x) = \begin{cases} H(x) & \text{für } |H(x)| \leq D, \\ D & \text{sonst,} \end{cases}$$

$$H'_D(x) = \begin{cases} H'(x) & \text{für } |H(x)| \leq D, \\ D & \text{sonst,} \end{cases}$$

über. Da H_D, H'_D stets beschränkt sind, können wir $H_D(A)$, $H'_D(B)$ jedenfalls bilden, und sie sind einander gewiß gleich. Die genannte Beschränktheitseigenschaft von $H(x)$ bzw. $H'(x)$ besagt nun, daß $H_C(x) = H_{C+1}(x)$ in einem $\mathfrak{M}$ nach II., 2., g) für A gilt, bzw. $H'_C(x) = H'_{C+1}(x)$

in einem $\mathfrak{M}$ nach II., 2., g) für B , d. h. (nach der Schlußbemerkung von II., 3.) $H_C(A) = H_{C+1}(A)$ bzw. $H'_C(B) = H'_{C+1}(B)$. Und da die linken Seiten einander gleich sind, und die rechten auch, ist beides gleichwertig, wie behauptet wurde.

Damit ist unsere Definition von $F(A_1, A_2, \dots)$ als einwandfrei erwiesen. Es ist noch zu erwähnen, daß sie der alten, soweit diese anwendbar ist, gleichwertig ist: wenn $F(x)$ eine Variable $x = x_1$ hat, und für diese ein Hermitesches A_1 einzusetzen ist, so können wir $A = A_1$ setzen, $G_1(x) = x$ ($G(A_1) = A_1 = A$), und das $F(A_1)$ im neuen Sinne ist $= F(A) = F(A_1)$ im alten Sinne, und auch die Beschränktheitsforderungen (d. h. die Bedingungen der Sinnvollheit) sind gleichlautend.

3. Wir verifizieren die Grundeigenschaften des neuen Funktionsbegriffes.

- a) Für $F(x_1, x_2, \dots) = \overline{G(x_1, x_2, \dots)}$ ist $F(A_1, A_2, \dots) = G(A_1^*, A_2^*, \dots)^*$.
- b) Für $F(x_1, x_2, \dots) = aG(x_1, x_2, \dots)$ ist $F(A_1, A_2, \dots) = aG(A_1, A_2, \dots)$.
- c) Für $F(x_1, x_2, \dots) = G(x_1, x_2, \dots) + H(x_1, x_2, \dots)$ ist $F(A_1, A_2, \dots) = G(A_1, A_2, \dots) + H(A_1, A_2, \dots)$.
- d) Für $F(x_1, x_2, \dots) = G(x_1, x_2, \dots) \cdot H(x_1, x_2, \dots)$ ist $F(A_1, A_2, \dots) = G(A_1, A_2, \dots) \cdot H(A_1, A_2, \dots)$.

(In a)–d) steht es mit der Sinnvollheit so: immer wenn die rechte Seite Sinn hat, hat auch die linke Sinn.)

- e) Für $F(x_1, x_2, \dots) = G(H_1(x_1, x_2, \dots), H_2(x_1, x_2, \dots), \dots)$ ist $F(A_1, A_2, \dots) = G(H_1(A_1, A_2, \dots), H_2(A_1, A_2, \dots), \dots)$.

(Nach a), d) sind nämlich alle $H_1(A_1, A_2, \dots)$, $H_2(A_1, A_2, \dots)$, $\dots$, wenn sie überhaupt Sinn haben, normal — in diesem Falle sind dann beide Seiten sinnvoll, oder keine.) Man erkennt, daß alle diese Eigenschaften mühelos durch Anwendung der Definition von $F(x_1, x_2, \dots)$ und Vergegenwärtigung der entsprechenden Eigenschaften des alten Funktionsbegriffes (vgl. II., 2., b)–f)) beweisbar sind.

Genaue Kriterien für die Gleichheit und Konvergenz beim neuen Funktionsbegriff anzugeben, so wie wir sie beim alten hatten (II., 2., g) und h), oder ausführlicher in der Schlußbemerkung von II., 3. und Satz 4, 5) ist nicht nötig: es genügt jeweils die definitorische Zurückführung des neuen Funktionsbegriffes auf den alten. Immerhin sei die folgende (roheste) Form des Konvergenzsatzes erwähnt:

- f) Es gelte überall $F_n(x_1, x_2, \dots) \rightarrow F(x_1, x_2, \dots)$, die $F_n(x_1, x_2, \dots)$ seien gleichmäßig beschränkt. (Dann ist auch $F(x_1, x_2, \dots)$ beschränkt, also alle $F_n(A_1, A_2, \dots)$, $F(A_1, A_2, \dots)$ sinnvoll.) Dann ist, im Sinne der starken Doppelkonvergenz in $\mathfrak{B}$, $F_n(A_1, A_2, \dots) \rightarrow F(A_1, A_2, \dots)$. Dies ist evident, da der alte Funktionsbegriff diese Eigenschaft auch hatte.

4. Wir beweisen nun die folgende Verschärfung der vorletzten Behauptung von III., 1:

Satz 8. $\mathbf{M}$ sei eine Abelsche Menge von Operatoren aus $\mathbf{B}$ (d. h. ihre Elemente seien linear-beschränkt, und für irgend zwei von ihnen, A, B , sei A mit B und mit B^* vertauschbar — also sind insbesondere alle normal). Dann gibt es einen Operator R derart, daß einerseits jedes A von $\mathbf{M}$ Funktion von R ist: $A = F_A(R)$, und andererseits eine endliche oder abzählbar unendliche Teilmenge $A_1, A_2, \dots$ von $\mathbf{M}$ existiert, so daß R Funktion von $A_1, A_2, \dots$ ist: $R = G(A_1, A_2, \dots)$.

Dabei kann R Hermitesch gewählt werden, und alle Funktionen an der Stelle 0 verschwindend: $F_A(0) = 0$ (für alle A von $\mathbf{M}$), $G(0, 0, \dots) = 0$.

Beweis: $\mathbf{R}(\mathbf{M})$ ist, wie $\mathbf{M}$ selbst, Abelsch, kann also (wie wir mehrfach erwähnten) als $\mathbf{R}(R)$, mit Hermitemischem R , geschrieben werden. Da also alle A von $\mathbf{M}$ zu $\mathbf{R}(\mathbf{M}) = \mathbf{R}(R)$ gehören, sind sie nach Satz 6 gleich $F_A(R)$, mit $F_A(0) = 0$. Umgekehrt gehört R zu $\mathbf{R}(R) = \mathbf{R}(\mathbf{M})$, ist also nach A, Seite 398, Satz 9, Limes einer stark doppelkonvergenten Folge solcher Ausdrücke, deren jeder durch endlich viele Operationen $A^*, aA, A+B, AB$ aus Elementen von $\mathbf{M}$ entsteht. Da in einen jeden dieser Ausdrücke nur endlich viele Elemente von $\mathbf{M}$ eingehen können, kommen für die gesamte Folge endlich oder abzählbar unendlich viele in Frage, wir nennen sie $A_1, A_2, \dots$. Jedes Element der Folge ist daher ein Polynom endlich vieler $A_1, A_2, \dots$ und $A_1^*, A_2^*, \dots$, ohne Absolutglied — wenn wir $A_1, A_2, \dots$ durch Variable $x_1, x_2, \dots$ und die $A_1^*, A_2^*, \dots$ durch deren Konjugierte $\bar{x}_1, \bar{x}_2, \dots$ ersetzen, so heiße das n -te $p_n(x_1, x_2, \dots)$. Nach Anm. ⁴²) ist $p_n(x_1, x_2, \dots)$ Φ -meßbar, ferner ist, da das Absolutglied fehlt, $p_n(0, 0, \dots) = 0$. Nach II., 2., α) sowie IV., 2., a)—d) ist dann das n -te Glied unserer Operatorenfolge $p_n(A_1, A_2, \dots)$. folglich haben wir $p_n(A_1, A_2, \dots) \rightarrow R$. Hieraus folgt $p_n(F_{A_1}(R), F_{A_2}(R), \dots) \rightarrow R$, aus Satz 5 folgt daher, daß es eine Teilfolge $n_1, n_2, \dots$ ($n_1 < n_2 < \dots$) von $n = 1, 2, \dots$ gibt, derart, daß $p_{n_\nu}(F_{A_1}(x), F_{A_2}(x), \dots) \rightarrow x$ (für $\nu \rightarrow \infty$) in einer x -Menge $\mathfrak{M}$ nach II., 2., g) (für R) stattfindet.

Sei nun

$$G(x_1, x_2, \dots) = \begin{cases} \lim_{\nu \rightarrow \infty} p_{n_\nu}(x_1, x_2, \dots), & \text{wenn dieser Limes existiert,} \\ 0 & \text{sonst.} \end{cases}$$

Nach Anm. ⁴²) ist dieses $G(x_1, x_2, \dots)$ Φ -meßbar, und nach dem vorhin Bemerkten gilt $G(F_{A_1}(x), F_{A_2}(x), \dots) = x$ in einer Menge $\mathfrak{M}$ nach II., 2., g) (für R), daher ist nach Definition $G(A_1, A_2, \dots) = R$. Aus $p_{n_\nu}(0, 0, \dots) = 0$ folgt $G(0, 0, \dots) = 0$.

Ferner gilt:

Satz 9. $\mathbf{M}$ sei wie in Satz 8. $\mathbf{R}(\mathbf{M})$ ist dann die Menge aller $F(A_1, A_2, \dots)$, wobei $F(x_1, x_2, \dots)$ alle Φ -meßbaren Funktionen mit $F(0, 0, \dots) = 0$ durchläuft, $A_1, A_2, \dots$ aber jede endliche oder abzählbar unendliche Teil-

menge von $\mathbf{M}$ sein darf. Das Resultat bleibt aber auch bestehen, wenn für $A_1, A_2, \dots$ nur eine einzige derartige Menge zugelassen wird, falls diese geeignet gewählt wird (was immer möglich ist).

Beweis: Seien $R, F_A(x), A_1, A_2, \dots, G(x_1, x_2, \dots)$ wie beim Beweise von Satz 8. Sind $B_1, B_2, \dots$ irgendwelche (endlich oder abzählbar unendlich viele) Elemente von $\mathbf{M}$, so ist jedes $F(B_1, B_2, \dots)$ wegen $B_1 = F_{B_1}(R), B_2 = F_{B_2}(R), \dots$ gleich $F(F_{B_1}(R), F_{B_2}(R), \dots) = H(R)$ mit $H(x) = F(F_{B_1}(x), F_{B_2}(x), \dots), H(0) = F(0, 0, \dots) = 0$, — also gehört es zu $\mathbf{R}(R) = \mathbf{R}(\mathbf{M})$. Umgekehrt ist jedes Element von $\mathbf{R}(\mathbf{M})$ eines von $\mathbf{R}(R)$, also gleich $J(R)$ mit $J(0) = 0$. Somit ist es gleich $J(G(A_1, A_2, \dots)) = K(A_1, A_2, \dots)$ mit $K(x_1, x_2, \dots) = J(G(x_1, x_2, \dots)), K(0) = J(0) = 0$. Somit hat es die gewünschte Form, sogar bei dieser einen speziellen Wahl der $B_1, B_2, \dots$ (nämlich gleich $A_1, A_2, \dots$).

Die in Satz 8, 9 vorkommenden Funktionen $G(x_1, x_2, \dots)$ bzw. $F(x_1, x_2, \dots)$ können, wie aus dem Beweise von Satz 8 hervorgeht, auf die Klasse derer eingeschränkt werden, die durch einen einzigen Grenzübergang im Sinne von Anm.⁴³⁾ aus Polynomen endlich vieler Variablen und ihrer konjugierten, $p(x_1, x_2, \dots)$, entstehen. Im Sinne von Anm.⁴²⁾ sind sie also von der ersten Baireschen Klasse, auch dann, wenn man die nullte auf die genannten $p(x_1, x_2, \dots)$ beschränkt. Definiert man dagegen die Baireschen Klassen durch überall konvergente Funktionenfolgen (vgl. IV., 2., f)), so braucht man, um diese Funktionen zu erreichen, drei sukzessive Grenzprozesse, wie man leicht erkennt. In diesem (allgemein üblichen) Sinne braucht man also alle Funktionen bis zur dritten Baireschen Klasse. Da also jede Funktion durch eine solche ersetzt werden kann, gilt dies auch für die $F_A(x)$ in Satz 8.

Schließlich machen wir noch eine, der Schlußbemerkung von III., 1. entsprechende, Feststellung. Wenn wir in der Aussage von Satz 9 die Bedingung $F(0, 0, \dots) = 0$ fallen lassen, so erhalten wir im Allgemeinen nicht $\mathbf{R}(\mathbf{M})$, sondern die Gesamtheit aller $a \cdot 1 + B$ (a eine komplexe Zahl, B ein Element von $\mathbf{R}(\mathbf{M})$), und dies ist nach A, Seite 397, Satz 7, die Menge $\mathbf{M}''$, das ist die Menge aller A , die mit allen C vertauschbar sind, welche mit B und B^* für alle B von $\mathbf{M}$ vertauschbar sind. Diese A sind also mit sämtlichen $F(A_1, A_2, \dots)$ (die $A_1, A_2, \dots$ nach Satz 9) identisch.

Anmerkungen.

¹⁾ Im Laufe dieser Abhandlung werden zwei frühere Arbeiten des Verfassers mehrfach zitiert werden, und zwar: „Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren“. Math. Ann., Bd. 102/1, Seite 49–131 (1929), und „Zur Algebra der Funktionaloperatoren und der Theorie der normalen Operatoren“, Math. Ann., Bd. 102/3, Seite 370–427 (1929). Wir werden sie mit E bzw. A bezeichnen. — Hier handelt es sich um A, insbesondere um den Schluß von Kap. III, § 2.

Es sei noch darauf hingewiesen, daß Funktionen von Operatoren bereits von F. Rieß betrachtet wurden, und u. a. bei seinem Beweise des Hilbertschen Satzes über die Spektralform beschränkter Hermitescher Operatoren eine wichtige Rolle spielten. Neuerdings kündigte M. Stone eine allgemeine Definition von Funktionen von Operatoren an (Proc. Nat. Ac., Bd. 16/2, Seite 172–175, 1930), ferner spielen besondere Operatorenfunktionen-Konstruktionen in zwei Arbeiten von A. Haar eine sehr bemerkenswerte Rolle (Math. Ztschr., Bd. 31/5, Seite 769–798, 1930, bzw. ebenda 1931 erscheinend).

²⁾ Die Hermiteschen (beschränkten, vgl. hierzu A. Einleitung, 1.) Operatoren A, B heißen vertauschbar, wenn $AB = BA$ gilt.

³⁾ Vgl. hierzu die Beschreibung des Hilbertschen Raumes in E, Kap. I (ferner Einleitung I–IV sowie Anhang I). Über Operatoren vgl. E, Kap. II, sowie Anhang II, § 1. Alles Notwendige ist übrigens in A, Einleitung, 1., nochmals zusammengestellt. Der gleich zu erwähnende Ring aller beschränkten Operatoren wird in A, Einleitung 2, eingeführt. — Punkte des Hilbertschen Raumes nennen wir $f, g, \dots, \varphi, \psi, \dots$, beschränkte Operatoren $A, B, \dots$, komplexe Zahlen $a, \dots$ (wie in E und A).

⁴⁾ A^* ist in der Matrizen-Terminologie die transponiert-konjugierte zu A . vgl. E, Anhang II, § 1 oder A, Ende von Einleitung, 1.

⁵⁾ Die Definition steht in A, Kap. II, § 1. Definition 1. Vgl. auch A. Kap. I, § 3, Zusatz zu Satz 5.

⁶⁾ Vgl. A, Kap. I, § 1, Definition 2.

⁷⁾ Vgl. A, Ende von Kap. III, § 2.

⁸⁾ D. h. wenn $A = A^*$ ist, bzw. A, A^* vertauschbar sind. Letzteres ist das Kriterium dafür, daß $\mathcal{R}(A)$ Abelsch ist, d. h. aus lauter miteinander vertauschbaren Operatoren besteht. Vgl. A, Kap. III, Anfang von § 1 und von § 2.

⁹⁾ Vgl. A, gegen Ende von Kap. II, § 1.

¹⁰⁾ Kap. III, § 2, Satz 10.

¹¹⁾ Vgl. W. H. Young, Proc. Cambridge Phil. Soc., Bd. 14. Seite 520–529, 1908; H. Lebesgue, Thèse, Ann. di Math., Bd. 7 (3), 1902.

¹²⁾ $g(u)$ braucht nicht mehr meßbar zu sein!

¹³⁾ Da $\psi(b)$ nach rechts halbstetig ist, nimmt es am linken Ende des Intervalles den im Inneren angenommenen Wert gewiß an, am rechten dagegen nur, wenn es dort stetig ist. D. h. wenn der Wert $\psi(b) = a$ nicht der Anfang eines Konstanzintervalles von $\varphi(a)$ ist. Das genannte Intervall ist also allenfalls inkl. linkes Ende zu verstehen, inkl. rechtes Ende aber nur im soeben genannten Falle.

¹⁴⁾ Dabei ist es wesentlich, daß $\mathfrak{P}_0$ alle I_n umfaßte, also $\mathfrak{P}_0$ auch: wäre $\mathfrak{P}_0$ zu einigen I_n fremd gewesen, so brauchte $\mathfrak{P}_0$ keins von beiden zu tun.

¹⁵⁾ Da $\Phi(a) = a$ gesetzt werden kann, folgt aus dieser Φ -Meßbarkeit die gewöhnliche Meßbarkeit.

¹⁶⁾ Für ein einziges $f_1(x)$, wobei dieses $= x$ ist (aber mit der zweiten Baireschen Klasse an Stelle der dritten), ist dies ein bekannter Vitalischer Satz. Das Wesentliche ist eben die Generalisation auf beliebige meßbare $f_1(x)$, $f_2(x)$, ... Zum Vitalischen Satz vgl. Rend. Lomb., Bd. 38, Seite 599 (1905).

¹⁷⁾ Sei etwa $H(x)$ von zweiter Klasse, also $H_1(x)$, $H_2(x)$, ... von erster, und $H_n(x) \rightarrow H(x)$. Wenn nun $H(x)$ in den Intervallen $I_1, I_2, \dots$ konstant gemacht werden soll, so erreichen wir dies, indem wir $H_1(x)$ nur in I_1 , $H_2(x)$ nur in I_1, I_2 , usw. konstant machen. — In unserem Falle ($H(x) = \bar{G}(x)$) und nach der Änderung $= \bar{G}(x)$) enthalten die I_n allenfalls ihre linken Enden, ob sie die rechten auch enthalten, ist ungewiß (vgl. Anm.¹³).

¹⁸⁾ Durch Zerlegung der Zahlengeraden in endlich viele Teile, deren jeder genau eine solche Stelle enthält, erkennt man, daß es genügt, sich mit einer einzigen solchen Stelle auseinanderzusetzen. Eine triviale Variablentransformation erlaubt noch, das Definitionsintervall zu $-1 < x < 1$, und die fragliche Stelle zu $x = 0$ zu normieren. Sei nunmehr die genannte Funktionenfolge $f_1(x)$, $f_2(x)$, ... Sei dann $\bar{f}_n(x) = f_n(x)$ für $|x| \geq \frac{1}{n}$ und für $x = 0$, in $-\frac{1}{n} \leq x \leq 0$ und in $0 \leq x \leq \frac{1}{n}$ aber linear. Dann sind die $\bar{f}_1(x)$, $\bar{f}_2(x)$, ... stetig, und haben dieselben Limites wie die $f_1(x)$, $f_2(x)$, ...

¹⁹⁾ $f_{m,n}(x) \rightarrow g_m(x)$ für $n \rightarrow \infty$, $g_m(x) \rightarrow \bar{G}(x)$ für $m \rightarrow \infty$.

²⁰⁾ Sei $f(x)$ nach rechts halbstetig, dann sind alle $\bar{f}_n(y) = n \int_y^{y+\frac{1}{n}} f(x) dx$ stetig, und $\bar{f}_n(x) \rightarrow f(x)$.

²¹⁾ Wenn $\mathfrak{M}$ alle im Definitionsintervalle von $\varphi(x)$ gelegenen Zahlen, ausschließlich seiner Konstanzintervalle, umfaßt, so ist dies erfüllt: denn es enthält dann alle Werte des $\psi(y)$. Übrigens erkennt man leicht, daß die y , für die $\psi(y)$ nicht in $\mathfrak{M}$ liegt, die folgenden sind: man nehme für jedes außerhalb von $\mathfrak{M}$ (aber im Definitionsintervalle von $\varphi(x)$) gelegene x die Zahl $y = \varphi(x)$, wenn x eine Stetigkeitsstelle von $\varphi(x)$ ist, und das ganze „Sprungintervall“ $\lim_{\substack{x' \rightarrow x \\ x' < x}} \varphi(x') \leq y \leq \varphi(x) = \lim_{\substack{x' \rightarrow x \\ x' > x}} \varphi(x')$, wenn x eine

Unstetigkeitsstelle ist. Daher ist von $\mathfrak{M}$ -s Komplementärmenge im Definitionsintervalle zu fordern: sie darf keine Unstetigkeitsstelle enthalten, und ihr durch $y = \varphi(x)$ vermitteltes Bild muß eine Lebesguesche Nullmenge sein.

²²⁾ $\mathfrak{M}$ soll Φ -meßbar sein, d. h. die Funktion, die in $\mathfrak{M} = 1$ und in der Komplementärmenge $= 0$ ist, soll es sein. Oder: für jedes monoton nichtfallende und nach rechts halbstetige $\varphi(x)$ soll die Menge der x mit $\varphi(x)$ aus $\mathfrak{M}$ im Lebesgueschen Sinne meßbar sein.

²³⁾ Man kann bekanntlich die obere und die untere Grenze des Realteils sowie des Imaginärteils aller $\varphi(x'_1) - \varphi(x'_1) + \dots + \varphi(x'_n) - \varphi(x'_n)$ mit $x'_1 < x'_1 < \dots < x'_n < x'_n < x$ (x fest, alles andere variabel) bilden, und diese bzw. gleich $\psi_1(x) - \Re \varphi(-\infty)$, $-\psi_2(x)$, $\psi_3(x) - \Im \varphi(-\infty)$, $-\psi_4(x)$ setzen. Das sind dann die gewünschten Funktionen.

²⁴⁾ Hilbert, Gött. Nachr., 1906, insbesondere Seite 189–190, Satz II. Unsere Bezeichnungen weichen etwas von den Hilbertschen ab, sie sind die in E und A verwendeten. Vgl. insbesondere E, Seite 91–92, Definition 17 und Satz 36, ferner Seite 114, Anhang II, Satz 9*.

²⁵⁾ Der Integrand, λ , ist ja stetig, und $(E(\lambda)f, g)$ von beschränkter Schwankung. Das letztere folgt daraus, daß

$$\Re(E(\lambda)f, g) = \left(E(\lambda) \frac{f+g}{2}, \frac{f+g}{2}\right) - \left(E(\lambda) \frac{f-g}{2}, \frac{f-g}{2}\right)$$

und $\Im(E(\lambda)f, g)$ entsprechend mit ig statt g ist, also $(E(\lambda)f, g)$ ein Aggregat von vier $(E(\lambda)h, h)$ mit den Koeffizienten $\pm 1, \pm i$; und $(E(\lambda)h, h) = |E(\lambda)h|^2$ ist in λ monoton nichtfallend und beschränkt (vgl. E, Seite 114).

²⁶⁾ Vgl. A, Seite 418, Anhang I.

²⁷⁾ Vgl. E, Seite 91, Definition 17. $E \leq F$ bedeutet $|Ef| \leq |Ff|$ für alle f , $E_n \rightarrow F$ bedeutet $E_n f \rightarrow Ff$ für alle f (im Sinne der „starken“ Konvergenz).

²⁸⁾ Da $\mathfrak{S}$ immer beschränkt ist, ist dies sicher der Fall, wenn $F(x)$ in jedem endlichen Intervalle beschränkt ist. $F(x) = x$ tut dies z. B.

²⁹⁾ Für $E \leq F$ gilt allgemein $|Ef|^2 - |Ff|^2 = |(E-F)f|^2$, vgl. E, Seite 78.

³⁰⁾ $E(\lambda_n) - E(\lambda_{n-1}) = E$ ist ein Projektionsoperator, und für diese gilt

$$|(Ef, g)| \leq |Ef| \cdot |Eg|$$

allgemein (vgl. E, Seite 93).

³¹⁾ Vgl. das Zitat und den Beweis in E, Seite 94, Anm. ⁶²⁾.

³²⁾ Für ein reelles $F(x)$ folgt also hieraus, daß $F(A)$ Hermitesch ist. Auf Grund des in 1. Gesagten genügt es sogar, wenn $F(x)$ im Spektrum von A , in $\mathfrak{S}$, reell ist.

³³⁾ Da $B = H(A)$ in $G(B)$ Hermitesch sein muß, und $y = H(x)$ in $G(y)$ reell, muß $H(x)$ stets reell sein (vgl. Anm. ³²⁾). Nach Satz 2 wird $F(x)$ $\mathcal{O}$ -meßbar sein, wenn $G(x)$. $H(x)$ es sind.

³⁴⁾ Genau wie in Anm. ²¹⁾ erkennt man, daß dies mit folgendem gleichbedeutend ist: die Komplementärmenge von $\mathfrak{M}$ (ihr allgemeines Element heiße ξ) darf, nachdem sie der Abbildung $x = \varrho^{-1}(\xi)$ unterworfen wurde, keine Unstetigkeitsstelle von $\mathcal{O}_f(x)$ enthalten, also sie selbst keine von $|E(\xi)f|^2$ (was für alle f der Fall ist, wenn sie zum Punktspektrum $\mathfrak{P}$ von A fremd ist) — ferner muß ihr durch $x = \varrho^{-1}(\xi)$ und $y = \mathcal{O}_f(x)$, d. h. ihr durch $y = |E(\xi)f|^2$ vermitteltes Bild muß eine Lebesguesche Nullmenge sein. Da diese Funktion in den Komplementärintervallen des Spektrums stets konstant ist, brauchen wir nur die im Spektrum $\mathfrak{S}$ von A liegenden Teile der Komplementärmenge von $\mathfrak{M}$ zu berücksichtigen (vgl. auch 1.). Wenn also $\mathfrak{M}$ ganz $\mathfrak{S}$ umfaßt, so ist alles erfüllt.

³⁵⁾ Vgl. A, Seite 383.

³⁶⁾ Vgl. E, Seite 65, Bedingung C.

³⁷⁾ Sei $F_n(A)f = f_n$, $F(A)f = \bar{f}$, dann ist zu zeigen: wenn „schwach“ $f_n \rightarrow \bar{f}$ ist, und $|f_n| \rightarrow |\bar{f}|$, so gilt $f_n \rightarrow \bar{f}$ „stark“ (zur Terminologie vgl. z. B. A, Seite 378–379). Es ist nämlich:

$|f_n - \bar{f}|^2 = (f_n - \bar{f}, f_n - \bar{f}) = |f_n|^2 + |\bar{f}|^2 - 2\Re(f_n, \bar{f}) \rightarrow |\bar{f}|^2 + |\bar{f}|^2 - 2\Re(\bar{f}, \bar{f}) = 0$, also wirklich „stark“ $f_n \rightarrow \bar{f}$.

³⁸⁾ Die Prämisse von 2., h) wird durch die Prämisse von γ) geliefert. Zwar erfolgt in 2., h) (bei gegebenem A und $F_n(x)$, $F(x)$) der Schluß von allen f auf alle f , aber man erkennt mühelos, daß ebensogut von einem gegebenen f (ohne Rücksicht auf die übrigen) auf dieses f geschlossen werden kann.

³⁹⁾ Er lautet so: Sei $\eta > 0$, $\mathfrak{M}_1, \mathfrak{M}_2, \dots, \mathfrak{N}$ nach Lebesgue meßbare Mengen, die Masse der $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ alle $\geq \eta$, das Maß von $\mathfrak{N}$ endlich, $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ alle Teilmengen von $\mathfrak{N}$. Dann hat die Menge aller Punkte, die unendlich vielen unter den $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ angehören, ein Maß $\geq \eta$. Vgl. z. B. de la Vallée Poussin, Cours d'Analyse infinitésimale, Ausg. 1/2 (Louvain-Paris, 1909), Seite 68–69.

⁴⁰⁾ Unter Voraussetzung der Richtigkeit der sog. Kontinuumshypothese (wonach jede abzählbare Zahlenmenge die Mächtigkeit des Kontinuums hat) läßt sich ein Beispiel

einer $\mathcal{O}$ -meßbaren Funktion mit einer Inversen, die nicht $\mathcal{O}$ -meßbar ist, angeben. Wir gehen hier nicht näher darauf ein. (Wird die $\mathcal{O}$ -Meßbarkeit durch gewöhnliche Meßbarkeit ersetzt, so existiert ein Gegenbeispiel gewiß.)

⁴¹⁾ Es ist leicht beweisbar, daß man sich ebensogut auf überall stetige $f_1(u), f_2(u), \dots$ beschränken könnte.

⁴²⁾ Als nullte Klasse definieren wir jetzt die Gesamtheit aller Funktionen, die nur von endlich vielen Variablen tatsächlich abhängen, und von diesen stetig. Den Grenzprozeß interpretieren wir im soeben angegebenen weiteren Sinne. So wird z. B. das durch

$$F(x_1, x_2, \dots) = \begin{cases} x_1 + x_2 + \dots & \text{wenn diese Reihe konvergiert} \\ 0 & \text{sonst} \end{cases}$$

definierte $F(x_1, x_2, \dots)$ zur ersten Klasse gehören.

⁴³⁾ Dies ist eine Art der Peanoschen Kurve, die von Steinhaus (Studia Mathematica, Bd. II, Seite 24–27, 1930) für solche Zwecke verwendet wurde. Sie beruht auf der dyadischen Entwicklung $z = 0, \zeta_1 \zeta_2 \dots = \sum_{\nu=1}^{\infty} \frac{\zeta_{\nu}}{2^{\nu}}$, $\zeta_{\nu} = 0, 1$, aber unendlich oft 0, die für alle Zahlen $0 \leq z < 1$ eindeutig möglich ist, und wird durch

$$y_1 = 0, \eta_1^{(1)} \eta_2^{(1)} \dots \quad \dots, \quad y_n = 0, \eta_1^{(n)} \eta_2^{(n)} \dots, \quad y = 0, \eta_1^{(1)} \dots \eta_1^{(n)} \eta_2^{(1)} \dots \eta_2^{(n)} \dots$$

definiert. Sie bildet den ganzen Würfel $0 \leq y_1 < 1, \dots, 0 \leq y_n < 1$ ein-eindeutig auf alle Punkte $y = 0, \xi_1 \xi_2 \dots$ des Intervalles $0 \leq y < 1$ ab, für die jede der Folgen $\xi_k, \xi_{k+n}, \xi_{k+2n}, \dots$ ($k = 1, \dots, n$) unendlich viele 0-en enthält — d. h. auf alle Punkte mit Ausnahme einer Lebesgueschen Nullmenge. Daß nun diese Abbildung die erwähnte Maß-erhaltende Eigenschaft hat, beweist man unschwer analog zu den Überlegungen, die Steinhaus a. a. O. anstellt.

⁴⁴⁾ Auch diese Abbildung läßt sich am einfachsten mit Hilfe der dyadischen Entwicklung angeben. Wenn die Zahl der $x_1, x_2, \dots$ endlich ist, etwa m , so sei

$$x_1 = 0, \xi_1^{(1)} \xi_2^{(1)} \dots, \quad \dots, \quad x_m = 0, \xi_1^{(m)} \xi_2^{(m)} \dots, \quad x = 0, \xi_1^{(1)} \dots \xi_1^{(m)} \xi_2^{(1)} \dots \xi_2^{(m)} \dots,$$

ist sie aber unendlich, so sei

$$x_1 = 0, \xi_1^{(1)} \xi_2^{(1)} \dots, \quad x_2 = 0, \xi_1^{(2)} \xi_2^{(2)} \dots \quad \dots, \quad x = 0, \xi_1^{(1)} \xi_1^{(2)} \xi_2^{(1)} \xi_2^{(2)} \xi_3^{(1)} \xi_3^{(2)} \xi_4^{(1)} \dots$$

Man verifiziert mühelos alle im Text formulierten Eigenschaften.

Algebraische Repräsentanten der Funktionen „bis auf eine Menge vom Maße Null“.

Einleitung.

Wir befassen uns mit komplexwertigen, beschränkten, für alle reellen x , und nur diese, definierten Funktionen $f(x)$, die im Lebesgueschen Sinne meßbar sind. (Verallgemeinerungen auf andere Definitionsbereiche sollen am Schlusse erfolgen.) Wenn zwei Funktionen dieser Art, $f(x)$ und $g(x)$, als äquivalent bezeichnet werden, falls die Menge aller x , für die $f(x) \neq g(x)$ ist, das Lebesguesche Maß 0 hat — in Symbolen $f \sim g$ — so erkennt man: $f \sim f$, aus $f \sim g$ folgt $g \sim f$, aus $f \sim g, g \sim h$ folgt $f \sim h$. Folglich zerfällt die Gesamtheit unserer f in paarweise elementfremde Äquivalenzklassen untereinander äquivalenter f .

Unter einem „allgemeineren Polynom“ $p(u_1, \dots, u_n)$ verstehen wir ein Linearaggregat beliebiger (aber endlich vieler) Potenzprodukte der $u_1, \dots, u_n$ und ihrer komplex-konjugierten $\bar{u}_1, \dots, \bar{u}_n$ (mit komplexen Zahlen als Koeffizienten). Da aus $f_1 \sim g_1, \dots, f_n \sim g_n$ $p(f_1, \dots, f_n) \sim p(g_1, \dots, g_n)$ folgt, hängt die Klasse $\mathfrak{R}$ von $p(f_1, \dots, f_n)$ nur von den Klassen $\mathfrak{R}_1, \dots, \mathfrak{R}_n$ der $f_1, \dots, f_n$ ab. Wir nennen $\mathfrak{R}$ darum $p(\mathfrak{R}_1, \dots, \mathfrak{R}_n)$ — insbesondere die Klasse der Konstanten a wieder a .

Im Zusammenhang mit gewissen Untersuchungen über unitäre Funktionaloperatoren, auf die hier nicht näher eingegangen werden soll, hat Herr A. Haar im mündlichen Gespräch mit dem Verfasser die folgende Frage formuliert: Ist es möglich jeder Äquivalenzklasse $\mathfrak{R}$ einen Repräsentanten $f = f_{\mathfrak{R}}$ (eine Funktion aus $\mathfrak{R}$!) derart zuzuordnen, daß für jedes allgemeinere Polynom $p(u_1, \dots, u_n)$ und irgendwelche Klassen $\mathfrak{R}_1, \dots, \mathfrak{R}_n$ aus $p(\mathfrak{R}_1, \dots, \mathfrak{R}_n) = 0$ identisch $p(f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)) \equiv 0$ folgt? (D. h.: Wenn für irgendwelche Funktionen $f_1(x), \dots, f_n(x)$ $p(f_1(x), \dots, f_n(x)) = 0$ mit Ausnahme einer x -Menge vom Lebesgueschen Maße 0 gilt, so gilt $p(f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)) = 0$ für alle x ohne Ausnahme — und dabei sollen $f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)$ nur von den bzw. Klassen $\mathfrak{R}_1, \dots, \mathfrak{R}_n$ von $f_1(x), \dots, f_n(x)$ abhängen, und auch selbst zu diesen Klassen gehören¹⁾).

¹⁾ Es ist erwähnenswert, daß diese Aufgabe ohne die Zusatzbedingung der Beschränktheit unlösbar ist. Sei nämlich ξ eine komplexe Zahl, und $f^*(x), g_\xi^*(x), e^*(x)$ die Repräsentanten der Klassen der Funktionen

$$f(x) = x, \quad g_\xi(x) = \begin{cases} \frac{1}{x - \xi} & \text{für } x \neq \xi \\ 0 & \text{für } x = \xi, \quad e(x) = 1. \end{cases}$$

Da für $x \neq \xi$ $(f(x) - \xi e(x)) g_\xi(x) - e(x) = 0$ gilt, muß für alle x $(f^*(x) - \xi e^*(x)) g_\xi^*(x) - e^*(x) = 0$ sein. Daher kann nie auf einmal $f^*(x) = \xi, e^*(x) = 1$ sein, insbesondere nicht $f^*(\xi) = \xi, e^*(\xi) = 1$.

Variieren wir nun ξ , dann müssen die Mengen der ξ mit $f^*(\xi) \neq f(\xi) (= \xi)$ und der ξ mit $e^*(\xi) \neq e(\xi) (= 1)$ zusammen alle ξ umfassen — was unmöglich ist, da beide das Lebesguesche Maß 0 haben sollten.

Ein solches Repräsentanten-System $f_{\mathfrak{R}}(x)$ wollen wir ein algebraisches Repräsentanten-System der Äquivalenzklassen nennen. Im ersten Teile der Arbeit zeigen wir, daß solche Systeme existieren, im zweiten nehmen wir einige Verallgemeinerungen vor, und beweisen einige allgemeingültige Eigenschaften dieser Systeme.

I. Beweis des Existenzsatzes.

1. Wir betrachten zuerst im Lebesgueschen Sinne meßbare Mengen reeller Zahlen. Zwei solche Mengen $\mathfrak{M}, \mathfrak{N}$ mögen äquivalent heißen, wenn ihre „Unterschiedsmenge“ $(\mathfrak{M} \dot{+} \mathfrak{N}) - (\mathfrak{M} \times \mathfrak{N})$ das Lebesguesche Maß 0 hat — in Symbolen $\mathfrak{M} \sim \mathfrak{N}$. Es ist $\mathfrak{M} \sim \mathfrak{N}$, aus $\mathfrak{M} \sim \mathfrak{N}$ folgt $\mathfrak{N} \sim \mathfrak{M}$, aus $\mathfrak{M} \sim \mathfrak{N}, \mathfrak{N} \sim \mathfrak{P}$ folgt $\mathfrak{M} \sim \mathfrak{P}$. Folglich zerfällt die Gesamtheit unserer $\mathfrak{M}$ in paarweise elementfremde Äquivalenzklassen untereinander äquivalenter $\mathfrak{M}$. Da aus $\mathfrak{M}_1 \sim \mathfrak{N}_1, \mathfrak{M}_2 \sim \mathfrak{N}_2$ $\mathfrak{M}_1 \dot{+} \mathfrak{M}_2 \sim \mathfrak{N}_1 \dot{+} \mathfrak{N}_2, \mathfrak{M}_1 \times \mathfrak{M}_2 \sim \mathfrak{N}_1 \times \mathfrak{N}_2, \mathfrak{M}_1 - \mathfrak{M}_2 \sim \mathfrak{N}_1 - \mathfrak{N}_2$ folgt, hängt die Klasse von $\mathfrak{M}_1 \dot{+} \mathfrak{M}_2$ bzw. $\mathfrak{M}_1 \times \mathfrak{M}_2$ bzw. $\mathfrak{M}_1 - \mathfrak{M}_2$ nur von den Klassen $\mathfrak{f}_1, \mathfrak{f}_2$ von $\mathfrak{M}_1, \mathfrak{M}_2$ ab. Wir nennen darum diese Klasse $\mathfrak{f}_1 \dot{+} \mathfrak{f}_2$ bzw. $\mathfrak{f}_1 \times \mathfrak{f}_2$ bzw. $\mathfrak{f}_1 - \mathfrak{f}_2$ — und die Klasse der leeren Menge 0 wieder 0, die der Menge 1 aller Zahlen wieder 1.

Wir wollen nun jeder Äquivalenzklasse $\mathfrak{f}$ einen Repräsentanten $\mathfrak{M} = \mathfrak{M}_{\mathfrak{f}}$ (eine Menge aus $\mathfrak{f}$!) derart zuordnen, daß jede Klassen-Relation $\mathfrak{f}_1 \times \cdots \times \mathfrak{f}_p \times (1 - \mathfrak{f}_1) \times \cdots \times (1 - \mathfrak{f}_q) = 0$ die Mengen-Relation $\mathfrak{M}_{\mathfrak{f}_1} \times \cdots \times \mathfrak{M}_{\mathfrak{f}_p} \times (1 - \mathfrak{M}_{\mathfrak{f}_1}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{f}_q}) = 0$ nach sich zieht. (D. h.: Wenn für irgendwelche Mengen $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{M}'_1, \dots, \mathfrak{M}'_q$ $\mathfrak{M}_1 \times \cdots \times \mathfrak{M}_p \times (1 - \mathfrak{M}'_1) \times \cdots \times (1 - \mathfrak{M}'_q)$ das Lebesguesche Maß Null hat, so ist $\mathfrak{M}_{\mathfrak{f}_1} \times \cdots \times \mathfrak{M}_{\mathfrak{f}_p} \times (1 - \mathfrak{M}_{\mathfrak{f}_1}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{f}_q})$ leer — und dabei sollen $\mathfrak{M}_{\mathfrak{f}_1}, \dots, \mathfrak{M}_{\mathfrak{f}_p}, \mathfrak{M}_{\mathfrak{f}'_1}, \dots, \mathfrak{M}_{\mathfrak{f}'_q}$ nur von den bzw. Klassen der $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{M}'_1, \dots, \mathfrak{M}'_q$ abhängen, und auch selbst zu diesen Klassen gehören.) —

Wenn $\mathfrak{M}$ eine meßbare Menge ist, und x eine reelle Zahl, so können wir für jedes $\delta > 0$ das Intervall $I_{\delta}(x) : x - \delta, x + \delta$ bilden, und den Bruch $\frac{\text{Maß } (I_{\delta}(x) \times \mathfrak{M})}{2\delta}$ betrachten. Für $\delta \rightarrow 0$ heiße sein Limes superior $\bar{\Delta}(x, \mathfrak{M})$, sein Limes inferior $\underline{\Delta}(x, \mathfrak{M})$, es ist stets $0 \leq \underline{\Delta} \leq \bar{\Delta} \leq 1$. Nach einem bekannten Satze von Lebesgue gilt in ganz $\mathfrak{M}$, mit Ausnahme einer Menge vom Maße 0, $\underline{\Delta}(x, \mathfrak{M}) = 1$ ²⁾; und da offenbar $\bar{\Delta}(x, \mathfrak{M}) = 1 - \underline{\Delta}(x, 1 - \mathfrak{M})$ ist, ergibt Anwendung auf $1 - \mathfrak{M}$: Außerhalb $\mathfrak{M}$ gilt, mit Ausnahme einer Menge vom Maße 0, $\bar{\Delta}(x, \mathfrak{M}) = 0$. Im ersteren Falle ist natürlich auch $\bar{\Delta} = 1$, im letzteren auch $\underline{\Delta} = 1$. Da aus $\mathfrak{M} \sim \mathfrak{N}$ $\bar{\Delta}(x, \mathfrak{M}) = \bar{\Delta}(x, \mathfrak{N}), \underline{\Delta}(x, \mathfrak{M}) = \underline{\Delta}(x, \mathfrak{N})$ folgt, hängen $\bar{\Delta}, \underline{\Delta}$ nur von der Klasse $\mathfrak{f}$ von $\mathfrak{M}$ ab, wir können also auch $\bar{\Delta}(x, \mathfrak{f}), \underline{\Delta}(x, \mathfrak{f})$ schreiben. Die Menge aller x mit $\bar{\Delta}(x, \mathfrak{f}) = \underline{\Delta}(x, \mathfrak{f}) = 1$ heiße $\bar{\mathfrak{M}}_{\mathfrak{f}}$, die Komplementärmenge aller x mit $\bar{\Delta}(x, \mathfrak{f}) = \underline{\Delta}(x, \mathfrak{f}) = 0$ heiße $\underline{\mathfrak{M}}_{\mathfrak{f}}$; nach dem vorhin Gesagten gehören $\bar{\mathfrak{M}}_{\mathfrak{f}}, \underline{\mathfrak{M}}_{\mathfrak{f}}$ beide zu $\mathfrak{f}$.

Man erkennt mühelos die folgenden Eigenschaften der $\bar{\mathfrak{M}}_{\mathfrak{f}}, \underline{\mathfrak{M}}_{\mathfrak{f}}$:

$$\bar{\mathfrak{M}}_0 = \underline{\mathfrak{M}}_0 = 0; \quad \bar{\mathfrak{M}}_{\mathfrak{f}} < \underline{\mathfrak{M}}_{\mathfrak{f}}; \quad \bar{\mathfrak{M}}_{1-\mathfrak{f}} = 1 - \underline{\mathfrak{M}}_{\mathfrak{f}},$$

$$\bar{\mathfrak{M}}_{1-\mathfrak{f}} = 1 - \underline{\mathfrak{M}}_{\mathfrak{f}}; \quad \bar{\mathfrak{M}}_{\mathfrak{f}} \dot{+} \underline{\mathfrak{M}}_{\mathfrak{f}} < \bar{\mathfrak{M}}_{\mathfrak{f} \dot{+} \mathfrak{f}} < \underline{\mathfrak{M}}_{\mathfrak{f}} \dot{+} \underline{\mathfrak{M}}_{\mathfrak{f}} = \underline{\mathfrak{M}}_{\mathfrak{f} \dot{+} \mathfrak{f}};$$

$$\bar{\mathfrak{M}}_{\mathfrak{f}_1 \times \mathfrak{f}_2} = \bar{\mathfrak{M}}_{\mathfrak{f}_1} \times \bar{\mathfrak{M}}_{\mathfrak{f}_2} < \bar{\mathfrak{M}}_{\mathfrak{f}_1 \times \mathfrak{f}_2} < \underline{\mathfrak{M}}_{\mathfrak{f}_1} \times \underline{\mathfrak{M}}_{\mathfrak{f}_2}. \quad -$$

Da für ein Repräsentanten-System $\mathfrak{M}_i$ u. a. $\mathfrak{f}_1 \times \mathfrak{f}_2 \times (1 - \mathfrak{f}_3) = 0$ $\mathfrak{M}_{\mathfrak{f}_1} \times \mathfrak{M}_{\mathfrak{f}_2} \times (1 - \mathfrak{M}_{\mathfrak{f}_3}) = 0$ nach sich zieht, d. h. $\mathfrak{M}_{\mathfrak{f}_1} \times \mathfrak{M}_{\mathfrak{f}_2} < \mathfrak{M}_{\mathfrak{f}_3}$; ferner $(1 - \mathfrak{f}_1) \times \mathfrak{f}_3 = 0$ $(1 - \mathfrak{M}_{\mathfrak{f}_1}) \times \mathfrak{M}_{\mathfrak{f}_3} = 0$ und $(1 - \mathfrak{f}_2) \times \mathfrak{f}_3 = 0$ $(1 - \mathfrak{M}_{\mathfrak{f}_2}) \times \mathfrak{M}_{\mathfrak{f}_3} = 0$, also $\mathfrak{M}_{\mathfrak{f}_1} > \mathfrak{M}_{\mathfrak{f}_3}$ und $\mathfrak{M}_{\mathfrak{f}_2} > \mathfrak{M}_{\mathfrak{f}_3}$, d. h. $\mathfrak{M}_{\mathfrak{f}_1} \times \mathfrak{M}_{\mathfrak{f}_2} > \mathfrak{M}_{\mathfrak{f}_3}$ — so haben alle drei Klassen-Relationen $\mathfrak{M}_{\mathfrak{f}_1} \times \mathfrak{M}_{\mathfrak{f}_2} = \mathfrak{M}_{\mathfrak{f}_3}$ zur Folge. Da sie für $\mathfrak{f}_1 \times \mathfrak{f}_2 = \mathfrak{f}_3$

²⁾ Vgl. Lebesgue, Annales de l'École Normale (3) 27 (1910), S. 407.

wirklich alle drei bestehen, ist $\mathfrak{M}_{\mathfrak{t}_1} \times \mathfrak{M}_{\mathfrak{t}_2} = \mathfrak{M}_{\mathfrak{t}_1 \times \mathfrak{t}_2}$. Durch analoge Betrachtungen zeigt man $\mathfrak{M}_{\mathfrak{t}_1} \dot{+} \mathfrak{M}_{\mathfrak{t}_2} = \mathfrak{M}_{\mathfrak{t}_1 \dot{+} \mathfrak{t}_2}$, $\mathfrak{M}_{1-\mathfrak{t}_1} = 1 - \mathfrak{M}_{\mathfrak{t}_1}$, sowie $\mathfrak{M}_0 = 0$, $\mathfrak{M}_1 = 1$. Da nun $\mathfrak{M}_{\mathfrak{t}_1}$ und $\mathfrak{M}_{\mathfrak{t}_2}$ nur diese Bedingungen nicht restlos erfüllen, sind sie als Repräsentanten unbrauchbar. Wir werden aber sehen, daß die Repräsentanten-Wahl mit $\mathfrak{M}_{\mathfrak{t}_1} < \mathfrak{M}_{\mathfrak{t}_2} < \mathfrak{M}_{\mathfrak{t}_1 \times \mathfrak{t}_2}$ gelingt.

Hierzu berufen wir uns auf den Wohlordnungssatz: Die Menge K aller Äquivalenzklassen $\mathfrak{f}$ sei nach einem geeigneten Wohlordnungstypus α_0 wohlgeordnet, d. h. die $\mathfrak{f}$ den Ordnungszahlen $\alpha < \alpha_0$ eineindeutig zugeordnet: $\mathfrak{f} = \mathfrak{f}_\alpha$ entspreche α . Wir wollen die $\mathfrak{M}_{\mathfrak{t}_\alpha}$ so bestimmen, daß allgemein

$$\mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \cdots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}) < \mathfrak{M}_{\mathfrak{t}_{\beta_1} \times \cdots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \cdots \times (1 - \mathfrak{t}_{\gamma_q})}$$

gilt. Dies hat die Konsequenz, daß die $\mathfrak{M}_{\mathfrak{t}_\alpha}$ eine Repräsentation bilden, wie wir es w. o. wünschten; denn für $\mathfrak{f}_{\beta_1} \times \cdots \times \mathfrak{f}_{\beta_p} \times (1 - \mathfrak{f}_{\gamma_1}) \times \cdots \times (1 - \mathfrak{f}_{\gamma_q}) = 0$ ist die rechte Seite $= \mathfrak{M}_0 = 0$, also auch die linke $= 0$. Ferner folgt daraus $\mathfrak{M}_{\mathfrak{t}_\alpha} < \mathfrak{M}_{\mathfrak{t}_\alpha}$. Ersetzen wir hierin $\mathfrak{f}_\alpha$ durch $1 - \mathfrak{f}_\alpha$ (das ja auch ein $\mathfrak{f}_{\alpha'}$ ist), so ergibt sich $\mathfrak{M}_{\mathfrak{t}_\alpha} > \mathfrak{M}_{\mathfrak{t}_\alpha}$ — also $\mathfrak{M}_{\mathfrak{t}_\alpha} < \mathfrak{M}_{\mathfrak{t}_\alpha} < \mathfrak{M}_{\mathfrak{t}_\alpha}$, d. h. das Repräsentanten-System ist von der vorher beschriebenen Art. Die Konstruktion dieser $\mathfrak{M}_{\mathfrak{t}_\alpha}$ vollziehen wir durch transfinite Rekursion: Nehmen wir an, daß für ein $\alpha_1 < \alpha_0$ die Definition der $\mathfrak{M}_{\mathfrak{t}_\alpha}$ mit $\alpha < \alpha_1$ bereits erfolgt sei, und daß die obige Mengen-Implikation für alle $\beta_1, \dots, \beta_p, \gamma_1, \dots, \gamma_q < \alpha_1$ erfüllt ist — es gilt dann $\mathfrak{M}_{\mathfrak{t}_{\alpha_1}}$ so zu definieren, daß diese Beziehung auch für alle $\beta_1, \dots, \beta_p, \gamma_1, \dots, \gamma_q \leq \alpha_1$ zutrifft. Da wir keine wirklich neue Bedingung vor uns haben, solange nicht ein β oder γ tatsächlich $= \alpha_1$ ist; da ferner, wenn sowohl ein β als auch ein $\gamma = \alpha_1$ ist, auf beiden Seiten unserer Relation 0 steht, also wieder keine Bedingung vorliegt; und da schließlich, wenn mehrere β (oder γ) $= \alpha_1$ sind, alle bis auf eines fortgelassen werden können, ohne unsere Beziehung zu ändern — so dürfen wir annehmen, daß genau ein $\beta = \alpha_1$ ist, und kein γ , oder umgekehrt. Sei dies β bzw. γ das letzte, d. h. β_p bzw. γ_q , schreiben wir noch $p + 1$ bzw. $q + 1$ für p bzw. q , so lauten die für $\mathfrak{M}_{\mathfrak{t}_{\alpha_1}}$ hinzukommenden Bedingungen so:

$$\left. \begin{aligned} & \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \cdots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}) \times \mathfrak{M}_{\mathfrak{t}_{\alpha_1}} \\ & < \mathfrak{M}_{\mathfrak{t}_{\beta_1} \times \cdots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \cdots \times (1 - \mathfrak{t}_{\gamma_q}) \times \mathfrak{t}_{\alpha_1}}, \\ & \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \cdots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}) \times (1 - \mathfrak{M}_{\mathfrak{t}_{\alpha_1}}) \\ & < \mathfrak{M}_{\mathfrak{t}_{\beta_1} \times \cdots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \cdots \times (1 - \mathfrak{t}_{\gamma_q}) \times \mathfrak{t}_{\alpha_1}}. \end{aligned} \right\} \beta_1, \dots, \beta_p, \gamma_1, \dots, \gamma_q < \alpha_1,$$

Diese Bedingungen können auch so geschrieben werden:

$$\begin{aligned} & \mathfrak{M}_{\mathfrak{t}_{\alpha_1}} < \mathfrak{M}_{\mathfrak{t}_{\beta_1} \times \cdots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \cdots \times (1 - \mathfrak{t}_{\gamma_q}) \times \mathfrak{t}_{\alpha_1}} \\ & + (1 - \mathfrak{M}_{\mathfrak{t}_{\beta_1}}) \times \cdots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}), \\ & \mathfrak{M}_{\mathfrak{t}_{\alpha_1}} > (1 - \mathfrak{M}_{\mathfrak{t}_{\beta_1} \times \cdots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \cdots \times (1 - \mathfrak{t}_{\gamma_q}) \times (1 - \mathfrak{t}_{\alpha_1}) \\ & \times \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \cdots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \cdots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}). \end{aligned}$$

Damit diese Bedingungen durch ein $\mathfrak{M}_{\mathfrak{t}_{\alpha_1}}$ für alle Komplexe $\beta_1, \dots, \beta_p, \gamma_1, \dots, \gamma_q < \alpha_1$ simultan befriedigt werden können, muß die rechte Seite der ersten Gleichung für jeden Komplex $\beta_1, \dots, \beta_p, \gamma_1, \dots, \gamma_q < \alpha_1$ > sein, als die rechte Seite der zweiten für jeden Komplex $\beta'_1, \dots, \beta'_p, \gamma'_1, \dots, \gamma'_q < \alpha_1$. (Dann können wir $\mathfrak{M}_{\mathfrak{t}_{\alpha_1}}$ z. B. der Vereinigungsmenge der rechten Seiten der zweiten Gleichung für alle zugelassenen β, γ -Komplexe gleichsetzen.) Die genannte Mengen-Implikation lautet:

$$\begin{aligned}
 & (1 - \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{t}_{\beta'_r} \times (1 - \mathfrak{t}_{\gamma'_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma'_s}) \times (1 - \mathfrak{t}_{\alpha_1})) \\
 & \times \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{M}_{\mathfrak{t}_{\beta'_r}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma'_1}}) \times \dots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma'_s}}) \\
 & < \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \dots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma_q}) \times \mathfrak{t}_{\alpha_1} \\
 & \dot{+} (1 - \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \dots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \dots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}})),
 \end{aligned}$$

oder

$$\begin{aligned}
 & \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{M}_{\mathfrak{t}_{\beta'_r}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma'_1}}) \times \dots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma'_s}}) \\
 & \times \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \dots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \dots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}) \\
 & < \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{t}_{\beta'_r} \times (1 - \mathfrak{t}_{\gamma'_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma'_s}) \times (1 - \mathfrak{t}_{\alpha_1}) \\
 & \dot{+} \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \dots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma_q}) \times \mathfrak{t}_{\alpha_1}^3).
 \end{aligned}$$

Nun ist die rechte Seite

$$\begin{aligned}
 & = \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{t}_{\beta'_r} \times (1 - \mathfrak{t}_{\gamma'_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma'_s}) \times (1 - \mathfrak{t}_{\alpha_1}) \dot{+} \mathfrak{t}_{\beta_1} \times \dots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma_q}) \times \mathfrak{t}_{\alpha_1} \\
 & > \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{t}_{\beta'_r} \times (1 - \mathfrak{t}_{\gamma'_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma'_s}) \times \mathfrak{t}_{\beta_1} \times \dots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma_q})^4).
 \end{aligned}$$

Somit genügt es, wenn

$$\begin{aligned}
 & \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{M}_{\mathfrak{t}_{\beta'_r}} \times \mathfrak{M}_{\mathfrak{t}_{\beta_1}} \times \dots \times \mathfrak{M}_{\mathfrak{t}_{\beta_p}} \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma'_1}}) \times \dots \\
 & \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma'_s}}) \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_1}}) \times \dots \times (1 - \mathfrak{M}_{\mathfrak{t}_{\gamma_q}}) \\
 & < \mathfrak{M}_{\mathfrak{t}_{\beta'_1}} \times \dots \times \mathfrak{t}_{\beta'_r} \times \mathfrak{t}_{\beta_1} \times \dots \times \mathfrak{t}_{\beta_p} \times (1 - \mathfrak{t}_{\gamma'_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma'_s}) \times (1 - \mathfrak{t}_{\gamma_1}) \times \dots \times (1 - \mathfrak{t}_{\gamma_q})
 \end{aligned}$$

gilt — aber dies folgt, da $\beta'_1, \dots, \beta'_r, \beta_1, \dots, \beta_p, \gamma_1, \dots, \gamma'_s, \gamma_1, \dots, \gamma_q < \alpha_1$ sind, aus unserer Prämisse.

Damit ist die Konstruktion des Repräsentanten-Systems der $\mathfrak{M}_{\mathfrak{t}_\alpha}, \alpha < \alpha_0$, d. h. der $\mathfrak{M}_{\mathfrak{t}}$, durchgeführt.

2. Wir betrachten nunmehr alle reellen, beschränkten, für alle reellen x , und nur diese, definierten Funktionen $f(x)$, die im Lebesgueschen Sinne meßbar sind. Unter Beachtung dieser Realitäts-Bedingung können wir die in der Einleitung beschriebene Äquivalenzklassen-Einteilung auch für diese Funktionen vornehmen, und ein algebraisches Repräsentanten-System $f_{\mathfrak{R}}(x)$ ihrer Klassen $\mathfrak{R}$ suchen — wobei wir an Stelle der komplexen allgemeineren Polynome sinngemäß die reellen gewöhnlichen Polynome treten lassen. D. h. wir beschränken die Fragestellung der Einleitung aufs Reelle. Diese Aufgabe werden wir nun, mit Hilfe des Mengen-Repräsentanten-Systems $\mathfrak{M}_{\mathfrak{t}}$, lösen.

a sei eine reelle rationale Zahl. $\mathfrak{R}(a, f)$ sei die Menge aller x mit $f(x) < a$, $\mathfrak{t}(a, f)$ die Klasse von $\mathfrak{R}(a, f)$, wir definieren $f^*(x)$ als untere Grenze aller rationalen a , für die x zu $\mathfrak{M}_{\mathfrak{t}(a, f)}$ gehört. Da $f(x)$ beschränkt ist, also für ein festes C stets $\leq C, \geq -C$, ist für $a > C$ bzw. $a \leq -C$ $\mathfrak{R}(a, f) = 1$ bzw. $= 0$, $\mathfrak{t}(a, f) = 1$ bzw. $= 0$, $\mathfrak{M}_{\mathfrak{t}(a, f)} = 1$ bzw. $= 0$, also x sicher Element bzw. sicher nicht Element von $\mathfrak{M}_{\mathfrak{t}(a, f)}$ — somit ist $f^*(x)$ endlich und $\leq C, \geq -C$, d. h. beschränkt. Für $f \sim g$ ist $\mathfrak{t}(a, f) = \mathfrak{t}(a, g)$ für jedes a , also $f^* \equiv g^*$, d. h. f^* hängt nur von der Klasse von f ab.

Die Unterschiedsmenge von $\mathfrak{R}(a, f)$ und $\mathfrak{M}_{\mathfrak{t}(a, f)}$ ist eine Menge vom Lebesgueschen Maße 0, also auch die Vereinigungsmenge derselben für alle rationalen a (bei festem f). Gehört aber x nicht zu ihr, so gehört es $\mathfrak{M}_{\mathfrak{t}(a, f)}$ dann und nur dann an, wenn es zu $\mathfrak{R}(a, f)$

³⁾ Die Mengen-Implikationen $(1 - \mathfrak{R}) \times \mathfrak{C} < \mathfrak{R}' \dot{+} (1 - \mathfrak{C}')$ und $\mathfrak{C} \times \mathfrak{C}' < \mathfrak{R} \dot{+} \mathfrak{R}'$ sind, wie man leicht verifiziert, gleichwertig.

⁴⁾ Wegen der leicht verifizierbaren Mengen-Implikation $\mathfrak{R} \times \mathfrak{C} \dot{+} \mathfrak{R}' \times (1 - \mathfrak{C}) > \mathfrak{R} \times \mathfrak{R}'$ ist

$$\mathfrak{M}_{\mathfrak{R} \times \mathfrak{C} \dot{+} \mathfrak{R}' \times (1 - \mathfrak{C})} > \mathfrak{M}_{\mathfrak{R} \times \mathfrak{R}'}$$

gehört, d. h. wenn $f(x) < a$ ist — also ist $f^*(x)$ dann die untere Grenze aller rationalen a mit $f(x) < a$, d. h. $f^*(x) = f(x)$. Die Menge aller x mit $f^*(x) \neq f(x)$ ist somit eine vom Lebesgueschen Maße 0, also $f^* \sim f$, d. h. f^* gehört zur Klasse von f .

Die Klasse von f sei $\mathfrak{R}$, wir nennen $f^* f_{\mathfrak{R}}$: $f_{\mathfrak{R}}$ hängt nur von $\mathfrak{R}$ ab, und gehört dazu. —

Sei $\varphi(u_1, \dots, u_n)$ eine (reelle) stetige Funktion, das in der Einleitung über allgemeinere Polynome Gesagte gilt wörtlich auch für $\varphi(u_1, \dots, u_n)$: Die Klasse von $\varphi(f_1, \dots, f_n)$ hängt nur von den bzw. Klassen $\mathfrak{R}_1, \dots, \mathfrak{R}_n$ der $f_1, \dots, f_n$ ab, und möge $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n)$ genannt werden. Wir nehmen nun $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n) = 0$ an, und wollen zeigen, daß für alle x $\varphi(f_{\mathfrak{R}_1}, \dots, f_{\mathfrak{R}_n}) = 0$ gilt.

Sei also per Absurdum $\varphi(f_{\mathfrak{R}_1}, \dots, f_{\mathfrak{R}_n}) \neq 0$ für ein $x = x_0$. Da $\varphi(u_1, \dots, u_n)$ stetig ist, existieren n Paare rationaler Zahlen $a_1, b_1, \dots, a_n, b_n$ mit $a_1 < f_{\mathfrak{R}_1}(x_0) < b_1, \dots, a_n < f_{\mathfrak{R}_n}(x_0) < b_n$, so daß für alle $a_1 \leq u_1 < b_1, \dots, a_n \leq u_n < b_n$ $\varphi(u_1, \dots, u_n) \neq 0$ ist. Also ist $\varphi(f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)) \neq 0$, wenn $a_1 \leq f_{\mathfrak{R}_1}(x) < b_1, \dots, a_n \leq f_{\mathfrak{R}_n}(x) < b_n$ ist. Da nun bis auf eine Menge vom Lebesgueschen Maße 0 $f_{\mathfrak{R}_1}(x) = f_1(x), \dots, f_{\mathfrak{R}_n}(x) = f_n(x)$ ist, ist bis auf eine solche x -Menge, $\varphi(f_1(x), \dots, f_n(x)) \neq 0$ Folge von $a_1 \leq f(x) < b_1, \dots, a_n \leq f_n(x) < b_n$. Und da wegen $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n) = 0$ mit Ausnahme einer eben solchen x -Menge $\varphi(f_1(x), \dots, f_n(x)) = 0$ ist, kann $a_1 \leq f_1(x) < b_1, \dots, a_n \leq f_n(x) < b_n$ nur in einer Menge vom Lebesgueschen Maße 0 gelten. Somit hat die Menge $\mathfrak{R}(b_1, f_1) \times (1 - \mathfrak{R}(a_1, f_1)) \times \dots \times \mathfrak{R}(b_n, f_n) \times (1 - \mathfrak{R}(a_n, f_n))$ das Maß 0, es ist $\mathfrak{I}(b_1, f_1) \times (1 - \mathfrak{I}(a_1, f_1)) \times \dots \times \mathfrak{I}(b_n, f_n) \times (1 - \mathfrak{I}(a_n, f_n)) = 0$, und daher $\mathfrak{M}_{\mathfrak{U}(b_1, f_1)} \times (1 - \mathfrak{M}_{\mathfrak{U}(a_1, f_1)}) \times \dots \times \mathfrak{M}_{\mathfrak{U}(b_n, f_n)} \times (1 - \mathfrak{M}_{\mathfrak{U}(a_n, f_n)}) = 0$. Nun gehört nach Definition von $f_{\mathfrak{R}_\nu}(x_0)$, und wegen $a_\nu < f_\nu(x_0) < b_\nu$, x_0 nicht zu $\mathfrak{M}_{\mathfrak{U}(a_\nu, f_\nu)}$, wohl aber zu $\mathfrak{M}_{\mathfrak{U}(b_\nu, f_\nu)}$ — also gehört es zu $\mathfrak{M}_{\mathfrak{U}(b_\nu, f_\nu)} \times (1 - \mathfrak{M}_{\mathfrak{U}(a_\nu, f_\nu)})$. Da dies für alle $\nu = 1, \dots, n$ gilt, gehört x_0 auch zum obenstehenden Durchschnitt dieser Mengen, obwohl derselbe leer sein sollte. Das ist unmöglich.

Somit bilden die $f_{\mathfrak{R}}$ tatsächlich ein Repräsentanten-System von der gewünschten Art.

3. Jetzt können wir uns den in der Einleitung beschriebenen komplexwertigen, beschränkten, für alle reellen x , und nur diese, definierten $f(x)$, die im Lebesgueschen Sinne meßbar sind, zuwenden. $\Re f(x), \Im f(x)$ sind reell, also Funktionen von der in 2. diskutierten Art, ihre Klassen hängen nur von der Klasse $\mathfrak{R}$ von $f(x)$ ab ($\Re u$ und $\Im u$ sind stetig, sogar allgemeinere Polynome: $\frac{1}{2}u + \frac{1}{2}\bar{u}$, $\frac{1}{2i}u - \frac{1}{2i}\bar{u}$), sie mögen $\Re \mathfrak{R}$ bzw. $\Im \mathfrak{R}$ genannt werden. Wir definieren nunmehr: $f_{\mathfrak{R}} = f_{\Re \mathfrak{R}} + i f_{\Im \mathfrak{R}}$. $f_{\mathfrak{R}}$ hängt nur von $\mathfrak{R}$ ab, und da bis auf eine x -Menge vom Lebesgueschen Maße 0 $f_{\Re \mathfrak{R}}(x) = \Re f(x), f_{\Im \mathfrak{R}}(x) = \Im f(x)$ gilt, also $f_{\mathfrak{R}}(x) = \Re f(x) + i \Im f(x) = f(x)$, gehört $f_{\mathfrak{R}}$ zu $\mathfrak{R}$.

Schließlich sei $\varphi(u_1, \dots, u_n)$ eine (für alle Komplexen $(u_1, \dots, u_n)$ stetige Funktion. Wir definieren die Klasse $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n)$ wie in der Einleitung und in 2. Seien $v_1, w_1, \dots, v_n, w_n$ $2n$ reelle Variable, wir setzen $\varphi(v_1, w_1, \dots, v_n, w_n) = \varphi(v_1 + iw_1, \dots, v_n + iw_n)$. Dann ist, wie man leicht verifiziert, $\varphi(\Re \mathfrak{R}_1, \Im \mathfrak{R}_1, \dots, \Re \mathfrak{R}_n, \Im \mathfrak{R}_n) = \varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n)$. Aus $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n) = 0$ folgt nunmehr $\varphi(\Re \mathfrak{R}_1, \Im \mathfrak{R}_1, \dots, \Re \mathfrak{R}_n, \Im \mathfrak{R}_n) = 0$, also nach 2. für alle x $\varphi(f_{\Re \mathfrak{R}_1}(x), f_{\Im \mathfrak{R}_1}(x), \dots, f_{\Re \mathfrak{R}_n}(x), f_{\Im \mathfrak{R}_n}(x)) = 0$, und daher nach Definition der $f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)$ $\varphi(f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)) = 0$.

Die $f_{\mathfrak{R}}(x)$ bilden also ein Repräsentanten-System von der gewünschten Art. Man

beachte, daß wir die charakteristische Eigenschaft derselben nicht nur für die allgemeineren Polynome $p(u_1, \dots, u_n)$ bewiesen haben, sondern für alle stetigen Funktionen $\varphi(u_1, \dots, u_n)$ überhaupt!

II. Eigenschaften der Repräsentanten-Systeme und Verallgemeinerungen.

1. In der Einleitung verlangten wir für die algebraischen Repräsentanten-Systeme der Äquivalenzklassen $f_{\mathfrak{R}}$ bloß für allgemeinere Polynome, daß aus $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n) = 0$ $\varphi(f_{\mathfrak{R}_1}, \dots, f_{\mathfrak{R}_n}) = 0$ (für alle x) folge — bei unserer Konstruktion erreichten wir jedoch, daß dies für alle stetigen φ galt. Dieser auffallende Umstand hat seine Erklärung darin, daß allgemein aus dem Erfülltsein unserer Bedingung für alle allgemeineren Polynome dasselbe für sämtliche stetigen Funktionen φ folgt, und dies soll jetzt bewiesen werden.

Sei nämlich $\varphi(\mathfrak{R}_1, \dots, \mathfrak{R}_n) = 0$, und trotzdem $\varphi(f_{\mathfrak{R}_1}(x_0), \dots, f_{\mathfrak{R}_n}(x_0)) \neq 0$. Da $\varphi(u_1, \dots, u_n)$ stetig ist, gibt es also ein $\delta > 0$, so daß für $|u_1 - f_{\mathfrak{R}_1}(x_0)|^2 + \dots + |u_n - f_{\mathfrak{R}_n}(x_0)|^2 < \delta^2$ auch noch $\varphi(u_1, \dots, u_n) \neq 0$ ist. Da mit Ausnahme einer Menge vom Maße 0 $\varphi(f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)) = 0$ ist, ist dann auch $|f_{\mathfrak{R}_1}(x) - f_{\mathfrak{R}_1}(x_0)|^2 + \dots + |f_{\mathfrak{R}_n}(x) - f_{\mathfrak{R}_n}(x_0)|^2 \geq \delta^2$. Wir setzen nun

$$g(x) = \text{Min} \left(\frac{1}{|f_{\mathfrak{R}_1}(x) - f_{\mathfrak{R}_1}(x_0)|^2 + \dots + |f_{\mathfrak{R}_n}(x) - f_{\mathfrak{R}_n}(x_0)|^2}, \frac{1}{\delta^2} \right),$$

dies ist eine beschränkte und meßbare Funktion. Nach obigem gilt, mit Ausnahme einer Menge vom Maße 0,

$$(|f_{\mathfrak{R}_1}(x) - f_{\mathfrak{R}_1}(x_0)|^2 + \dots + |f_{\mathfrak{R}_n}(x) - f_{\mathfrak{R}_n}(x_0)|^2) g(x) - 1 = 0.$$

Nun ist $(|u_1 - f_{\mathfrak{R}_1}(x_0)|^2 + \dots + |u_n - f_{\mathfrak{R}_n}(x_0)|^2) u_{n+1} - 1$ ein allgemeineres Polynom in $u_1, \dots, u_n, u_{n+1}$, und dieses ist für $u_1 = f_{\mathfrak{R}_1}(x), \dots, u_n = f_{\mathfrak{R}_n}(x), u_{n+1} = g(x)$, mit Ausnahme einer Menge vom Maße 0, gleich 0. $f_{\mathfrak{R}_1}(x), \dots, f_{\mathfrak{R}_n}(x)$ haben die Klassen $\mathfrak{R}_1, \dots, \mathfrak{R}_n$, $g(x)$ habe die Klasse $\mathfrak{Q}$, dann gilt für die Repräsentanten

$$(|f_{\mathfrak{R}_1}(x) - f_{\mathfrak{R}_1}(x_0)|^2 + \dots + |f_{\mathfrak{R}_n}(x) - f_{\mathfrak{R}_n}(x_0)|^2) f_{\mathfrak{Q}}(x) - 1 = 0$$

für alle x . Trotzdem ist linke Seite für $x = x_0$ offenkundig $= -1 \neq 0$. Dies ist ein Widerspruch.

2. An Stelle des Lebesgueschen Maßes, daß allen unseren bisherigen Betrachtungen zugrunde lag, können wir auch andere, verwandte Begriffsbildungen treten lassen.

So können wir z. B. eine für alle reellen x , und nur diese, definierte Funktion $\varphi(x)$ einführen, welche monoton nicht-abnehmend ist, und an die Stelle der Mengen vom Maße 0 diejenigen Mengen treten lassen, über welche erstreckt das Lebesgue-Stieltjesche Integral $\int d\varphi(x) = 0$ ist.

Unter Verwendung einer auf Lebesgue zurückgehenden Definition des Lebesgue-Stieltjeschen Integrals, können wir dies auch so machen. Sei $\psi(y)$ die Inverse von $\varphi(x)$, die, um auch streckenweise konstante oder stellenweise unstetige $\varphi(x)$ (wie sie z. B. in der Operatoretheorie vorkommen) mit zu erfassen, etwa so definiert wird: $\psi(y)$ ist die obere Grenze aller x , für die $\varphi(x) \leq y$ ist (findet dies nie oder immer statt, so sei $\psi(y) = -\infty$ bzw. $= +\infty$ — oder auch undefiniert). Das über eine Menge $\mathfrak{M}$ erstreckte Integral $\int d\varphi(x)$ ist dann gleich dem Maß der Menge $\mathfrak{M}_\varphi$ aller y , für die $\psi(y)$ zu $\mathfrak{M}$ gehört. Wir lassen also an Stelle der Mengen $\mathfrak{M}$ vom Maße 0 diejenigen $\mathfrak{M}$ treten, für welche $\mathfrak{M}_\varphi$ das Maß 0 hat.

Ferner betrachten wir, an Stelle der beschränkten, im Lebesgueschen Sinne meßbaren Funktionen $f(x)$, diejenigen beschränkten $f(x)$, für die $f(\varphi(x))$ im Lebesgueschen

Sinne meßbar ist. Und unter $f \sim g$ verstehen wir, daß für die Menge $\mathfrak{M}$ aller x mit $f(x) \neq g(x)$ $\mathfrak{M}_\varphi$ das Maß 0 hat. Somit bedeutet $f \sim g$ $f(\varphi) \sim g(\varphi)$. Für diesen neuen Äquivalenzbegriff $\sim$ können wir wiederum Äquivalenzklassen bilden, und unser ursprüngliches Problem erneut formulieren.

Da jedoch das neue Problem für die f mit dem alten für die $f(\varphi)$ gleichwertig ist ⁵⁾, existieren auch jetzt die Repräsentanten. —

Eine andere Verallgemeinerungs-Möglichkeit wäre die, die f nicht von einer, sondern von mehreren reellen Variablen abhängen zu lassen. (Hierher gehört auch der Fall, daß die Variable x komplex ist.) Da es eindeutige und maßgleiche Abbildungen des n -dimensionalen Raumes auf die Zahlengerade gibt ⁶⁾, gelten unsere Resultate ⁴⁾ auch dann; übrigens gelten unsere Überlegungen ohnehin beinahe ohne jede Änderung auch im Raume.

⁵⁾ Allerdings brauchen die Repräsentanten, die nach unserem ursprünglichen Verfahren gewonnen werden, nicht in den Konstanzintervallen von φ konstant zu sein, was die $f(\varphi)$ gerade charakterisiert. Wir hätten uns aber bei unserem ursprünglichen Verfahren konsequent auf solche Funktionen beschränken können, die in den Konstanzintervallen von φ konstant sind (bzw. auf solche Mengen, die von jedem dieser Intervalle keinen oder alle Punkte enthalten), und wären genau so ans Ziel gekommen.

⁶⁾ Vgl. z. B. a. a. O. ²⁾.

ERRATA for
Die Eindeutigkeit der Schrödingerschen Operatoren

p. 228, line 15 should read:

$$U(\gamma)f_{\alpha, \beta} = e^{i\beta\gamma} f_{\alpha+\gamma, \beta}, \quad V(\delta)f_{\alpha, \beta} = e^{-i\alpha\delta} f_{\alpha, \beta+\delta}$$

Die Eindeutigkeit der Schrödingerschen Operatoren

1. Die sogenannte Vertauschungsrelation

$$PQ - QP = \frac{h}{2\pi i} 1$$

ist in der neuen Quantentheorie von fundamentaler Bedeutung, sie ist es, die den „Koordinaten-Operator“ R und den „Impuls-Operator“ P im wesentlichen definiert¹⁾. Mathematisch gesprochen, liegt darin die folgende Annahme: Seien P, Q zwei Hermitesche Funktionaloperatoren des Hilbertschen Raumes, dann werden sie durch die Vertauschungsrelation bis auf eine Drehung des Hilbertschen Raumes, d. i. eine unitäre Transformation U , eindeutig festgelegt²⁾. Es liegt im Wesen der Sache, daß noch der Zusatz gemacht werden muß: vorausgesetzt, daß P, Q ein irreduzibles System bilden (vgl. weiter unten Anm. ⁶⁾). Wird nun, wie es sich durch die Schrödingersche Fassung der Quantentheorie als besonders günstig erwies, der Hilbertsche Raum als Funktionenraum interpretiert — der Einfachheit halber etwa als Raum aller komplexen Funktionen $f(q)$ ($-\infty < q < +\infty$)

mit endlichem $\int_{-\infty}^{+\infty} |f(q)|^2 dq$ —, so gibt es nach Schrödinger ein besonders einfaches Lösungssystem der Vertauschungsrelation

$$Q: f(q) \rightarrow q f(q), \quad P: f(q) \rightarrow \frac{h}{2\pi i} \frac{d}{dq} f(q) \text{ } ^3).$$

¹⁾ Vgl. Born-Heisenberg-Jordan, Zeitschr. f. Phys. **34** (1925), S. 858—888, ferner Dirac, Proc. Roy. Soc. **109** (1925) u. f. Besonders in der letztgenannten Darstellung ist die Rolle dieser Relation fundamental. Einen interessanten Versuch zur Begründung des im folgenden zu diskutierenden Eindeutigkeitssatzes machte Jordan, Zeitschr. f. Phys. **37** (1926), S. 383—386. Indessen beruht dieser auf Konvergenzannahmen über Potenzreihen unbeschränkter Operatoren, deren Gültigkeitsbereich fraglich ist.

²⁾ Dieselbe bewirkt ein Ersetzen von P, Q durch UPU^{-1}, UQU^{-1} , wodurch weder der Hermitesche Charakter noch das Bestehen der Vertauschungsrelation berührt wird.

³⁾ Vgl. Schrödinger, Annalen d. Phys. **79** (1926), S. 734—756.

Sind nun dies die im wesentlichen einzigen (irreduziblen) Lösungen der Vertauschungsrelation?

Indessen ist die Aufgabe in dieser Form nicht genügend präzise formuliert. Denn als P, Q sind, wie es die Schrödingerschen Lösungen zeigen, auch unbeschränkte, nicht überall definierte Operatoren ins Auge zu fassen, und für diese wird der Operator $PQ - QP$ nicht überall definiert sein, während es der (auf der anderen Seite der Vertauschungsrelation stehende) Operator $\frac{h}{2\pi i} 1$ ist. Die beiden Seiten können also nur gleichgesetzt werden, wenn ihre Definitionsbereiche (d. h. der der linken Seite) näher umschrieben werden. Dieser Schwierigkeit kann man folgendermaßen aus dem Wege gehen:

Durch formale Operatorenrechnung folgt aus der Vertauschungsrelation ($F(x)$ analytisch, $F'(x)$ seine Ableitung, vgl. Anm. 1))

$$PF(Q) - F(Q)P = \frac{h}{2\pi i} F'(Q),$$

und hieraus für $F(x) = e^{\frac{2\pi i}{h}\beta x}$

$$e^{-\frac{2\pi i}{h}\beta Q} P e^{\frac{2\pi i}{h}\beta Q} = P + \beta 1.$$

Hieraus folgt wieder formal

$$e^{-\frac{2\pi i}{h}\beta Q} F(P) e^{\frac{2\pi i}{h}\beta Q} = F(P + \beta 1),$$

und somit für $F(x) = e^{\frac{2\pi i}{h}\alpha x}$

$$e^{\frac{2\pi i}{h}\alpha P} e^{\frac{2\pi i}{h}\beta Q} = e^{\frac{2\pi i}{h}\alpha\beta} \cdot e^{\frac{2\pi i}{h}\beta Q} e^{\frac{2\pi i}{h}\alpha P}.$$

Diese Gleichung ist von Weyl aufgestellt und als Ersatz der Vertauschungsrelation vorgeschlagen worden⁴⁾. Ihr großer Vorzug besteht in folgendem: Es ist unter Umständen möglich, mit Hilfe der Operatoren P, Q einparametrische Scharen $U(\alpha) = e^{\frac{2\pi i}{h}\alpha P}$, $V(\beta) = e^{\frac{2\pi i}{h}\beta Q}$ zu definieren, die unitär sind, und dem Multiplikationsgesetz

$$U(\alpha)U(\beta) = U(\alpha + \beta), \quad V(\alpha)V(\beta) = V(\alpha + \beta)$$

genügen⁵⁾. Dann stehen auf beiden Seiten der Weylschen Gleichung

$$U(\alpha)V(\beta) = e^{\frac{2\pi i}{h}\alpha\beta} \cdot V(\beta)U(\alpha)$$

⁴⁾ Vgl. Weyl, Zeitschr. f. Phys. 46 (1928), Seite 1—46.

⁵⁾ Vgl. Weyl, Anm. 4), ferner Stone, Proc. of Nat. Academy 1930. Im Schrödingerschen Falle wird, wie man leicht erkennt:

$$U(\alpha): f(q) \rightarrow f(q + \alpha), \quad V(\beta): f(q) \rightarrow e^{\frac{2\pi i}{h}\beta q} f(q).$$

unitäre, also beschränkte und überall definierte Operatoren, so daß ihr Sinn ein völlig klarer ist.

Es bliebe daher zu zeigen, daß die einzigen irreduziblen Lösungen⁶⁾ der Weylschen Gleichungen die Schrödingerschen (d. h. die aus Anm. 5)) sind. Beweisansätze hierfür gab Stone (vgl. Anm. 5)) an, jedoch ist bisher ein Beweis auf dieser Grundlage, wie mir Herr Stone freundlichst mitteilte, nicht erbracht worden.

Im folgenden soll der genannte Eindeutigkeitssatz bewiesen werden. Wir werden sogar alle (auch die reduziblen) Lösungen angeben können.

2. Sei $\mathfrak{H}$ der Hilbertsche Raum (etwa durch alle Folgen komplexer Zahlen $\{x_1, x_2, \dots\}$ mit endlichem $\sum_{n=1}^{\infty} |x_n|^2$ realisiert; oder auch durch alle komplexen Funktionen $f(q)$, $-\infty < q < +\infty$, mit endlichem $\int_{-\infty}^{+\infty} |f(q)|^2 dq$).

Wir benützen die geometrische Terminologie in $\mathfrak{H}$, indem wir das „innere Produkt“ (f, g) (gleich $\sum_{n=1}^{\infty} x_n \bar{y}_n$ bzw. $\int_{-\infty}^{+\infty} f(q) \bar{g(q)} dq$) und den „absoluten Betrag“ $|f| = \sqrt{(f, f)}$ (gleich $\sqrt{\sum_{n=1}^{\infty} |x_n|^2}$ bzw. $\sqrt{\int_{-\infty}^{+\infty} |f(q)|^2 dq}$) einführen⁷⁾.

Wir werden ausschließlich beschränkt-lineare (überall definierte) Operatoren in $\mathfrak{H}$ betrachten, den transponiert-konjugierten Operator des Operators A nennen wir A^* (er ist durch $(Af, g) = (f, A^*g)$, $(f, Ag) = (A^*f, g)$ definiert). Wir erwähnen noch eins: Wenn der Operator $A(\alpha)$ vom Parameter α abhängt, so nennen wir diese Abhängigkeit meßbar, wenn alle Funktionen $(A(\alpha)f, g)$ (dies sind komplexe Zahlenfunktionen der reellen Zahlenvariablen α , dagegen betrachten wir f, g als Parameter) im Lebesgueschen Sinne in α meßbar sind⁸⁾. Daß mit $A(\alpha)$ auch $\alpha A(\alpha)$, $A(\alpha)^*$ und mit $A(\alpha), B(\alpha)$ auch $A(\alpha) + B(\alpha)$ meßbar ist, ist klar, aber auch $A(\alpha)B(\alpha)$ ist es. Dies folgt aus den bekannten Regeln der Matrizenmultiplikation, oder auch direkt, $\varphi_1, \varphi_2, \dots$ sei ein vollständiges, normiertes

⁶⁾ Ein System von Operatoren $A, B, \dots$ (im vorliegenden Falle besteht es aus allen $U(\alpha)$ und $V(\beta)$) heißt irreduzibel, wenn es außer O und dem vollen Hilbertschen Raume keine abgeschlossene Linearmannigfaltigkeit (d. h. Hyperebene) $\mathfrak{M}$ mit der folgenden Eigenschaft gibt: mit f gehören auch $Af, Bf, \dots$ zu $\mathfrak{M}$. Vgl. auch die Ausführungen im Buch von Born und Jordan, Elementare Quantenmechanik. Berlin 1930.

⁷⁾ Vgl. E. Schmidt, Rend. Circ. Mat. Palermo 25 (1908), S. 57–73, ferner die Arbeit des Verf., Math. Annalen 102 (1930), S. 49–131, an die die Bezeichnungsweise anlehnt.

⁸⁾ Sei $\varphi_1, \varphi_2, \dots$ ein vollständiges, normiertes Orthogonalsystem in $\mathfrak{H}$. Dann ist $f = \sum_{n=1}^{\infty} x_n \varphi_n$, $g = \sum_{n=1}^{\infty} y_n \varphi_n$, also $(A(\alpha)f, g) = \lim_{M, N \rightarrow \infty} \sum_{m=1}^M \sum_{n=1}^N x_m \bar{y}_n (A(\alpha)\varphi_m, \varphi_n)$. Somit genügt die Meßbarkeit der $(A(\alpha)\varphi_m, \varphi_n)$, d. h. der Matrizenelemente von $A(\alpha)$ im Koordinatensystem der $\varphi_1, \varphi_2, \dots$.

Orthogonalsystem:

$$\begin{aligned}(A(\alpha)B(\alpha)f, g) &= (B(\alpha)f, A(\alpha)^*g) = \sum_{n=1}^{\infty} (B(\alpha)f, \varphi_n)(\varphi_n, A^*(\alpha)g) \\ &= \sum_{n=1}^{\infty} \overline{(A(\alpha)g, \varphi_n)} (B(\alpha)f, \varphi_n).\end{aligned}$$

Dasselbe gilt, wenn an der Stelle von α mehrere Variable $\alpha, \beta, \dots$ stehen. Wir kehren nun zu unserem Problem zurück, ersetzen aber in $V(\beta)$ β durch $\frac{h}{2\pi}\beta$. Dann lautet es so:

Alle $U(\alpha), V(\beta)$ seien unitäre Operatoren, die meßbar von α, β abhängen. Es gelten die Relationen

$$\begin{aligned}U(\alpha)U(\beta) &= U(\alpha + \beta), & V(\alpha)V(\beta) &= V(\alpha + \beta), \\ U(\alpha)V(\beta) &= e^{i\alpha\beta}V(\beta)U(\alpha).\end{aligned}$$

Alle derartigen Systeme sind zu bestimmen.

Wenn wir die (von α, β meßbar abhängende, unitäre) Operatorenschar

$$S(\alpha, \beta) = e^{-\frac{1}{2}i\alpha\beta}U(\alpha)V(\beta) = e^{\frac{1}{2}i\alpha\beta}V(\beta)U(\alpha)$$

einführen, so können wir die obigen Relationen zu

$$S(\alpha, \beta)S(\gamma, \delta) = e^{\frac{1}{2}i(\alpha\delta - \beta\gamma)}S(\alpha + \gamma, \beta + \delta)$$

zusammenfassen. Infolgedessen ist $S(0, 0)$ die Einheit, und daher $S(-\alpha, -\beta)$ zu $S(\alpha, \beta)$ reziprok, also $S(\alpha, \beta)^* = S(-\alpha, -\beta)$. Es sollen nun Linearaggregate der $S(\alpha, \beta)$ betrachtet werden, diese werden folgendermaßen definiert: Sei $\alpha(\alpha, \beta)$ eine über die ganze α, β -Ebene absolut integrierbare Funktion, dann ist wegen der Schwarzschen Ungleichheit

$$|(S(\alpha, \beta)f, g)| \leq |S(\alpha, \beta)f| \cdot |g| = |f| \cdot |g|,$$

d. h. beschränkt, also auch das Integral

$$\iint \alpha(\alpha, \beta) (S(\alpha, \beta)f, g) d\alpha d\beta$$

absolut konvergent. Und zwar ist es, wenn wir $c = \iint |\alpha(\alpha, \beta)| d\alpha d\beta$ setzen, absolut $\leq c \cdot |f| \cdot |g|$. Dabei ist es in f linear und in g konjugiert-linear. Daher ist ein Satz von F. Rieß anwendbar⁹⁾, wonach bei festem f ein f^* existiert, so daß dieser Ausdruck für jedes $g = (f^*, g)$ ist, und zwar ist $|f^*| \leq c \cdot |f|$. f^* ist durch f bestimmt, und zwar ist die Abhängigkeit linear, wir können also einen linearen Operator A durch $Af = f^*$ definieren, nach der obigen Formel ist A auch beschränkt. Wir schreiben symbolisch

$$A = \iint \alpha(\alpha, \beta) S(\alpha, \beta) d\alpha d\beta,$$

obwohl die Definition eigentlich

$$(Af, g) = \iint \alpha(\alpha, \beta) (S(\alpha, \beta)f, g) d\alpha d\beta$$

lautet. $\alpha(\alpha, \beta)$ heiße der Kern von A .

⁹⁾ Vgl. auch a. a. O. Anm. ⁷⁾, Math. Annalen 102 (1930), S. 94, Anm. ⁵³⁾.

Wir beweisen einige Rechenregeln für diese Operatoren. Daß aA den Kern $a a(\alpha, \beta)$ hat, ist klar, A^* hat wegen $S(\alpha, \beta)^* = S(-\alpha, -\beta)$ den Kern $\overline{a}(-\alpha, -\beta)$, $AS(u, v)$ und $S(u, v)A$ wegen der Multiplikationsregel der $S(\alpha, \beta)$ den Kern

$$e^{\frac{1}{2}i(\alpha v - \beta u)} a(\alpha - u, \beta - v) \quad \text{bzw.} \quad e^{-\frac{1}{2}i(\alpha v - \beta u)} a(\alpha - u, \beta - v).$$

Haben A, B die bzw. Kerne $a(\alpha, \beta), b(\alpha, \beta)$, so hat $A + B$ offenbar $a(\alpha, \beta) + b(\alpha, \beta)$, bei AB dagegen ist eine kleine Rechnung notwendig:

$$\begin{aligned} (ABf, g) &= (Bf, A^*g) = \iint b(\alpha, \beta) (S(\alpha, \beta)f, A^*g) d\alpha d\beta \\ &= \iint b(\alpha, \beta) (AS(\alpha, \beta)f, g) d\alpha d\beta \\ &= \iiint b(\alpha, \beta) e^{\frac{1}{2}i(\gamma\beta - \delta\alpha)} a(\gamma - \alpha, \delta - \beta) (S(\gamma, \delta)f, g) d\alpha d\beta d\gamma d\delta \\ &= \iint \left[\iint e^{\frac{1}{2}i(\gamma\beta - \delta\alpha)} a(\gamma - \alpha, \delta - \beta) b(\alpha, \beta) d\alpha d\beta \right] (S(\gamma, \delta)f, g) d\gamma d\delta. \end{aligned}$$

Der Kern von AB ist also (statt γ, δ schreiben wir wieder α, β , statt α, β ξ, η) $\iint e^{\frac{1}{2}i(\alpha\eta - \beta\xi)} a(\alpha - \xi, \beta - \eta) b(\xi, \eta) d\xi d\eta$. (Die absolute Integrierbarkeit folgt aus der Deduktion.)

Schließlich zeigen wir: wenn A verschwindet, so ist auch sein Kern (bis auf eine Lebesguesche Nullmenge) gleich 0. Aus $A = 0$ folgt nämlich $S(-u, -v)AS(u, v) = 0$, also, da dieses den Kern $e^{i(\alpha v - \beta u)} a(\alpha, \beta)$ hat,

$$\iint e^{i(\alpha v - \beta u)} a(\alpha, \beta) (S(\alpha, \beta)f, g) d\alpha d\beta = 0.$$

Somit ist jedenfalls

$$\iint P(\alpha, \beta) a(\alpha, \beta) (S(\alpha, \beta)f, g) d\alpha d\beta = 0,$$

wenn $P(\alpha, \beta)$ ein Linearaggregat von endlich vielen $e^{i(k\alpha + l\beta)}$ ist, also für jedes trigonometrische Polynom mit einer Periode $p > 0$ in α, β . Da der zweite Faktor absolut integrierbar ist, und der dritte beschränkt, können wir mit dem ersten ($P(\alpha, \beta)$) Grenzübergänge ausführen, falls dieser dabei gleichmäßig beschränkt bleibt. So können wir die Klasse der $P(\alpha, \beta)$ sukzessiv erweitern: 1. zu allen stetigen Funktionen mit einer Periode $p > 0$ in α, β , 2. zu allen beschränkten stetigen Funktionen, 3. zu allen beschränkten Funktionen der ersten Baireschen Klasse. Wenn also $\mathfrak{R}$ ein beliebiges (endliches) Rechteck in der α, β -Ebene ist, so können wir $P(\alpha, \beta)$ in $\mathfrak{R}$ gleich 1 und außerhalb $= 0$ setzen, es wird:

$$\iint_{\mathfrak{R}} a(\alpha, \beta) (S(\alpha, \beta)f, g) d\alpha d\beta = 0$$

für alle diese $\mathfrak{R}$. Daher ist (mit Ausnahme einer α, β -Nullmenge) $a(\alpha, \beta) (S(\alpha, \beta)f, g) = 0$. Dies gilt bei festem f, g , ist aber nur f fest, während g ein vollständiges normiertes Orthogonalsystem durchläuft, so gilt es für dieses f und alle genannten g auch noch mit Ausnahme einer

α, β -Nullmenge. In diesem Falle ist aber $\alpha(\alpha, \beta) S(\alpha, \beta) f = 0$. Da nun für $f \neq 0$ $|S(\alpha, \beta) f| = |f| > 0$ ist, muß dann $\alpha(\alpha, \beta) = 0$ sein — womit alles bewiesen ist.

3. Die Lösung des Eindeutigkeits-Problems wird durch Betrachten des Operators

$$A = \iint e^{-\frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2} S(\alpha, \beta) d\alpha d\beta$$

gewonnen. Dieser Operator ist nach unseren bisherigen Resultaten hermitesch (d. h. $A = A^*$) und $\neq 0$, und hat außerdem die bemerkenswerte Eigenschaft, daß sich $AS(u, v)A$ nur um einen Zahlenfaktor von A unterscheidet. In der Tat hat A den Kern $e^{-\frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2}$, also $S(u, v)A$ den Kern $e^{-\frac{1}{2}i(\alpha v - \beta u)} e^{-\frac{1}{4}(\alpha - u)^2 - \frac{1}{4}(\beta - v)^2}$ also $AS(u, v)A$

$$\begin{aligned} & \iint e^{\frac{1}{2}i(\alpha\eta - \beta\xi)} e^{-\frac{1}{4}(\alpha - \xi)^2 - \frac{1}{4}(\beta - \eta)^2} e^{-\frac{1}{2}i(\xi v - \eta u)} e^{-\frac{1}{4}(\xi - u)^2 - \frac{1}{4}(\eta - v)^2} d\xi d\eta \\ &= \iint e^{-\frac{1}{2}\xi^2 - \frac{1}{2}\eta^2 + (\frac{1}{2}\alpha + \frac{1}{2}u - \frac{1}{2}i\beta - \frac{1}{2}iv)\xi + (\frac{1}{2}\beta + \frac{1}{2}v + \frac{1}{2}i\alpha + \frac{1}{2}iu)\eta - \frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2 - \frac{1}{4}u^2 - \frac{1}{4}v^2} d\xi d\eta \\ &= e^{-\frac{1}{4}u^2 - \frac{1}{4}v^2} e^{-\frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2} \iint e^{-\frac{1}{2}\xi^2 + \frac{1}{2}(\alpha + u - i\beta - iv)\xi - \frac{1}{2}\eta^2 + \frac{1}{2}(\beta + v + i\alpha + iu)\eta} d\xi d\eta \\ &= e^{-\frac{1}{4}u^2 - \frac{1}{4}v^2} e^{-\frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2} \iint e^{-\frac{1}{2}(\xi - \frac{1}{2}(\alpha + u - i\beta - iv))^2 - \frac{1}{2}(\eta - \frac{1}{2}(\beta + v + i\alpha + iu))^2} d\xi d\eta \\ &= e^{-\frac{1}{4}u^2 - \frac{1}{4}v^2} e^{-\frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2} \iint e^{-\frac{1}{2}x^2 - \frac{1}{2}y^2} dx dy \\ &= 2\pi e^{-\frac{1}{4}u^2 - \frac{1}{4}v^2} \cdot e^{-\frac{1}{4}\alpha^2 - \frac{1}{4}\beta^2}. \end{aligned}$$

Hierin ist der erste Faktor konstant und der zweite der Kern von A , also gilt

$$AS(u, v)A = 2\pi e^{-\frac{1}{4}u^2 - \frac{1}{4}v^2} \cdot A.$$

Wir betrachten nun die Lösungen der Gleichung $Af = 2\pi f$, da A linear-beschränkt ist, bilden sie eine abgeschlossene Linearmannigfaltigkeit im Hilbertschen Raume, $\mathfrak{M}$. Jede von ihnen hat die Form Ag (mit $g = \frac{1}{2\pi}f$), und umgekehrt gehört jedes Ag zu ihnen, da $A^2 = 2\pi \cdot A$ ist (man setze in der obigen Gleichung $u = v = 0$). Die zu $\mathfrak{M}$ orthogonale abgeschlossene Linearmannigfaltigkeit sei $\mathfrak{N}$, die Elemente f von $\mathfrak{N}$ sind durch Orthogonalität zu allen Elementen von $\mathfrak{M}$ gekennzeichnet, d. h. zu allen Ag . D. h.: immer $(f, Ag) = 0$, oder: immer $(Af, g) = 0$, oder: $Af = 0$.

Wenn f, g zu $\mathfrak{M}$ gehören, so ist

$$\begin{aligned} (S(\alpha, \beta)f, S(\gamma, \delta)g) &= \frac{1}{4\pi^2} (S(\alpha, \beta)Af, S(\gamma, \delta)Ag) \\ &= \frac{1}{4\pi^2} (AS(-\gamma, -\delta)S(\alpha, \beta)Af, g) = \frac{1}{4\pi^2} e^{\frac{1}{2}i(\alpha\delta - \beta\gamma)} (AS(\alpha - \gamma, \beta - \delta)Af, g) \\ &= \frac{1}{2\pi} e^{-\frac{1}{4}(\alpha - \gamma)^2 - \frac{1}{4}(\beta - \delta)^2 + \frac{1}{2}i(\alpha\delta - \beta\gamma)} (Af, g) = e^{-\frac{1}{4}(\alpha - \gamma)^2 - \frac{1}{4}(\beta - \delta)^2 + \frac{1}{2}i(\alpha\delta - \beta\gamma)} (f, g). \end{aligned}$$

Sei nun $\varphi_1, \varphi_2, \dots$ ein normiertes Orthogonalsystem, welches in $\mathfrak{M}$ vollständig ist (d. h. $\mathfrak{M}$ aufgespannt, die Zahl seiner Elemente ist endlich oder abzählbar unendlich). Aus $(\varphi_m, \varphi_n) = \delta_{mn}$ (d. i. 1 für $m = n$, 0 sonst) folgt

$$(S(\alpha, \beta) \varphi_n, S(\gamma, \delta) \varphi_n) = e^{-\frac{1}{4}(\alpha - \gamma)^2 - \frac{1}{4}(\beta - \delta)^2 + \frac{1}{2}i(\alpha\delta - \beta\gamma)} \delta_{nn}.$$

Die durch alle $S(\alpha, \beta) \varphi_n$ (n fest, α, β variieren) aufgespannte abgeschlossene Linearmannigfaltigkeit heie $\mathfrak{P}_n$, nach der obigen Formel sind für $m \neq n$ $\mathfrak{P}_m, \mathfrak{P}_n$ zueinander orthogonal. Die $\mathfrak{P}_1, \mathfrak{P}_2, \dots$ mögen zusammen die abgeschlossene Linearmannigfaltigkeit $\mathfrak{S}$ aufspannen, die zu $\mathfrak{S}$ komplementäre abgeschlossene Linearmannigfaltigkeit sei $\mathfrak{T}$.

Da jedes $S(\gamma, \delta)$ die $S(\alpha, \beta) \varphi_n$ (bis auf Zahlenfaktoren) in ebensolche transformiert, bildet es $\mathfrak{P}_n$ auf ein Teil von sich ab; da dasselbe für $S(\gamma, \delta)^{-1} = S(-\gamma, -\delta)$ gilt, ist $\mathfrak{P}_n$ genau invariant. Also ist auch $\mathfrak{S}$ und $\mathfrak{T}$ invariant. $\mathfrak{S}$ umfat alle $\mathfrak{P}_n$, also alle φ_n , also $\mathfrak{M}$, daher liegt $\mathfrak{T}$ in $\mathfrak{N}$. Somit gilt in $\mathfrak{T}$ stets $Af = 0$. Nun gelten alle unsere über A angestellten Betrachtungen schon in $\mathfrak{T}$, denn die $S(\alpha, \beta)$ können als Operatoren in $\mathfrak{T}$ angesehen werden, da dieses ihnen gegenüber invariant ist. Da nun in $\mathfrak{T}$ $A \equiv 0$ ist, kann nach unserem Beweise in $\mathfrak{T}$ niemals $f \neq 0$ sein. Also enthält $\mathfrak{T}$ nur die 0, $\mathfrak{S}$ ist der Hilbertsche Raum, d. h.: $\mathfrak{P}_1, \mathfrak{P}_2, \dots$ spannen den vollen Hilbertschen Raum auf.

Der Hilbertsche Raum erscheint somit als in eine endliche oder abzählbar unendliche Zahl von (paarweise orthogonalen) Unterräumen $\mathfrak{P}_1, \mathfrak{P}_2, \dots$ zerlegt; jeder derselben ist gegenüber allen $S(\alpha, \beta)$ invariant, es genügt also das Verhalten der $S(\alpha, \beta)$ (d. i. der $U(\alpha), V(\beta)$) in einem jeden derselben gesondert zu ermitteln, um über sie restlos informiert zu sein. (Im Falle der Irreduzibilität darf es natürlich nur ein $\mathfrak{P}_n$ geben, und dieses ist dann der volle Hilbertsche Raum.) In $\mathfrak{P}_n$ wissen wir nun über die $S(\alpha, \beta)$ die folgenden Tatsachen:

Wir nennen $\mathfrak{P}_n$ $\mathfrak{P}$, $S(\alpha, \beta) \varphi_n$ $f_{\alpha, \beta}$. Dann gilt:

$$S(\gamma, \delta) f_{\alpha, \beta} = e^{\frac{1}{2}i(\beta\gamma - \alpha\delta)} f_{\alpha + \gamma, \beta + \delta},$$

$$(f_{\alpha, \beta}, f_{\gamma, \delta}) = e^{-\frac{1}{4}(\alpha - \gamma)^2 - \frac{1}{4}(\beta - \delta)^2 + \frac{1}{2}i(\alpha\delta - \beta\gamma)},$$

und die Linearaggregate endlich vieler $f_{\alpha, \beta}$ (die beliebig wählbar sind!) liegen in $\mathfrak{P}$ überall dicht.

Wenn wir nun zeigen können, daß irgend zwei solche $\mathfrak{P}$, in deren jedem eine Schar von unitären Operatoren $S(\alpha, \beta)$ und Punkten $f_{\alpha, \beta}$ mit den obigen Eigenschaften gegeben ist, isomorph sind, so sind wir am Ziele. Isomorphie bedeutet: Existenz einer ein-eindeutigen, linearen und längentreuen Abbildung der beiden $\mathfrak{P}$ aufeinander, die die $f_{\alpha, \beta}$ und $S(\alpha, \beta)$ des einen $\mathfrak{P}$ in dieselben des anderen überführt.

Unsere Formel für $(f_{\alpha\beta}, f_{\gamma\delta})$ erlaubt für jedes Linearaggregat endlich vieler $f_{\alpha\beta}$ den Absolutwert zu berechnen, wenn wir also die gleichlautenden $f_{\alpha\beta}$ -Linearaggregate beider $\mathfrak{P}$ einander zuordnen, so haben wir eine eindeutige, lineare und längentreue Abbildung dieser Mengen aufeinander. Da sie in den bzw. $\mathfrak{P}$ überall dicht sind, sind sie stetig auf die ganzen $\mathfrak{P}$ ausdehnbar. Dabei bleiben Linearität und Längentreue, also auch Eindeutigkeit, erhalten. Die bzw. $f_{\alpha,\beta}$ entsprechen einander. Wegen der $S(\gamma, \delta) f_{\alpha,\beta}$ -Formeln gehen auch die $S(\gamma, \delta)$ in ihre entsprechenden über: wenigstens für die $f_{\alpha,\beta}$, aber dann auch für deren Linearaggregate und Häufungspunkte — also in ganz $\mathfrak{P}$. Damit ist, wenn wir wieder zu den $U(\alpha), V(\beta)$ zurückkehren, folgendes bewiesen:

Ein System unitärer Operatoren $U(\alpha), V(\beta)$ nebst einem System von Punkten $f_{\alpha,\beta}$, die zusammen den ganzen Hilbertschen Raum aufspannen, ist durch die Eigenschaften

$$V(\gamma) f_{\alpha,\beta} = e^{\frac{1}{2}i\beta\gamma} f_{\alpha+\gamma,\beta}, \quad V(\delta) f_{\alpha,\beta} = e^{-\frac{1}{2}i\alpha\delta} f_{\alpha,\beta+\delta},$$

$$(f_{\alpha,\beta}, f_{\gamma,\delta}) = e^{-\frac{1}{4}(\alpha-\gamma)^2 - \frac{1}{4}(\beta-\delta)^2 + \frac{1}{2}i(\alpha\delta - \beta\gamma)},$$

bis auf eine unitäre Transformation¹⁰⁾ eindeutig festgelegt.

Ein System unitärer Operatoren $U(\alpha), V(\beta)$ mit den Weylschen Multiplikations-Relationen (vgl. 1.) ist entweder eines der soeben genannten Systeme, oder es entsteht dadurch, daß der Hilbertsche Raum in endlich oder abzählbar unendlich viele (paarweise orthogonale, Hilbertsche) Unterräume zerfällt, und in jedem derselben ein solches System angenommen wird. D. h. es entsteht durch das Zusammenfügen derselben.

Die irreduziblen Lösungen sind offenbar die ersteren (die Zahl der linear unabhängigen Lösungen von $Af = 2\pi \cdot f$ nimmt für sie ihren Minimalwert 1 an).

4. Zum Schluß noch einige Zusatzbemerkungen.

Im Falle der Schrödingerschen Operatoren haben wir $U(\alpha), V(\beta)$ in Anm. 5) angegeben: $f(q) \rightarrow f(q+\alpha)$, $f(q) \rightarrow e^{i\beta q} f(q)$ ($\frac{h}{2\pi}\beta$ statt $\beta!$), daher ist $S(\alpha, \beta)$: $f(q) \rightarrow e^{i\beta(q+\frac{\alpha}{2})} f(q+\alpha)$, und, wie man leicht berechnet, A : $f(q) \rightarrow \sqrt[4]{4\pi} \int_{-\infty}^{+\infty} e^{-\frac{1}{2}q^2 - \frac{1}{2}q'^2} f(q') dq'$. Die einzigen Lösungen von $Af = 2\pi \cdot f$ sind also die $c \cdot e^{-\frac{1}{2}q^2}$, somit ist $\varphi_1 = \varphi_1(q) \equiv \pi^{-\frac{1}{4}} e^{-\frac{1}{2}q^2}$, und

$$f_{\alpha,\beta} = f_{\alpha,\beta}(q) \equiv \pi^{-\frac{1}{4}} e^{-\frac{1}{2}(q+\alpha)^2 + i\beta(q+\frac{\alpha}{2})} \equiv \pi^{-\frac{1}{4}} e^{-\frac{1}{2}q^2 + (-\alpha + i\beta)q + (-\frac{\alpha^2}{2} + \frac{i\alpha\beta}{2}}.$$

Man verifiziert dann leicht unsere diesbezüglichen Formeln. —

¹⁰⁾ Sie heie U , dann bewirkt sie (vgl. Anm. 2))

$$U(\alpha) \rightarrow U U(\alpha) U^{-1}, \quad V(\beta) \rightarrow U V(\beta) U^{-1}, \quad f_{\alpha,\beta} \rightarrow U f_{\alpha,\beta}.$$

Bei quantenmechanischen Problemen mit k ($= 1, 2, \dots$) Freiheitsgraden tritt das allgemeine Vertauschungsrelationen-System

$$\left. \begin{aligned} P_m P_n &= P_n P_m, & Q_m Q_n &= Q_n Q_m, \\ P_m Q_n - Q_n P_m &= \frac{\hbar}{2\pi i} \delta_{mn} \cdot 1 \end{aligned} \right\} \quad (m, n = 1, \dots, k)$$

auf. Durch Einführen von

$$U_n(\alpha) = e^{\frac{2\pi i}{\hbar} \alpha P_n}, \quad V_n(\beta) = e^{i\beta Q_n} \quad (n = 1, \dots, k)$$

entstehen hieraus die Weylschen Relationen

$$\begin{aligned} U_n(\alpha) U_n(\beta) &= U_n(\alpha + \beta), & V_m(\alpha) V_n(\beta) &= V_n(\alpha + \beta), \\ U_n(\alpha) U_m(\beta) &= U_m(\beta) U_n(\alpha), & V_m(\alpha) V_n(\beta) &= V_n(\beta) V_m(\alpha), \\ U_m(\alpha) V_n(\beta) &= e^{i\delta_{mn} \cdot \alpha \beta} \cdot V_n(\beta) U_m(\alpha). \end{aligned}$$

Auch hier ist der Eindeutigkeitsbeweis mit unseren Methoden aus 2. bis 3. durchführbar. Wir setzen

$$\begin{aligned} S(\alpha_1, \dots, \alpha_k, \beta_1, \dots, \beta_k) &= e^{-\frac{1}{2}i(\alpha_1\beta_1 + \dots + \alpha_k\beta_k)} U(\alpha_1) \dots U(\alpha_k) V(\beta_1) \dots V(\beta_k) \\ &= e^{\frac{1}{2}i(\alpha_1\beta_1 + \dots + \alpha_k\beta_k)} V(\beta_1) \dots V(\beta_k) U(\alpha_1) \dots U(\alpha_k) \end{aligned}$$

und

$$\iint \dots \iint e^{-\frac{1}{4}\alpha_1^2 - \dots - \frac{1}{4}\alpha_k^2 - \frac{1}{4}\beta_1^2 - \dots - \frac{1}{4}\beta_k^2} S(\alpha_1, \dots, \alpha_k, \beta_1, \dots, \beta_k) d\alpha_1 \dots d\alpha_k d\beta_1 \dots$$

Dann können wir, genau wie in 3.,

$$A S(u_1, \dots, u_k, v_1, \dots, v_k) A = (2\pi)^k e^{-\frac{1}{4}u_1^2 - \dots - \frac{1}{4}u_k^2 - \frac{1}{4}v_1^2 - \dots - \frac{1}{4}v_k^2} A$$

beweisen, und (durch Untersuchen der Lösungen von $Af = (2\pi)^k \cdot f$) genau wie dort ans Ziel kommen. —

Vom allgemeinen darstellungstheoretischen Gesichtspunkte aus betrachtet, ist unsere Betrachtungsweise mit der Frobeniusschen Behandlung endlicher Gruppen mittels ihrer „charakteristischen Einheiten“ verwandt bzw. mit der Weylschen Untersuchung abgeschlossener kontinuierlicher Gruppen mit Hilfe ihrer Gruppennzahlen¹¹⁾. Die Operatoren $A = \iint \alpha(\alpha, \beta) S(\alpha, \beta) d\alpha d\beta$ sind nämlich als „Gruppennzahlen“ der $S(\alpha, \beta)$ -Gruppe deutbar, und $AS(u, v)A = c_{u,v}A$ ($c_{u,v}$ eine Zahl!) ist der definierenden Eigenschaft der „primitiven“ charakteristischen Einheiten gleichwertig.

¹¹⁾ Vgl. Frobenius, Berl. Ber. 1896 u. f., Peter und Weyl, Math. Annalen 96 (1926), S. 737—755.

Bemerkungen zu den Ausführungen von Herrn St. Leśniewski über meine Arbeit »Zur Hilbertschen Beweistheorie«.

1. In seiner Abhandlung „Grundzüge eines neuen Systems der Grundlagen der Mathematik“ (Fund. Math. XIV, S. 1—81) hat sich Herr Leśniewski u. a. mit den Versuchen anderer Autoren, die Mathematik und Logik für den Zweck von „Grundlagen-Untersuchungen“ in ein rein formales System von Zeichen und Operationsvorschriften mit diesen Zeichen zu übersetzen, kritisch auseinandergesetzt. Insbesondere hat er das von mir („Zur Hilbertschen Beweistheorie“, Math. Zeitschr. 26, S. 1—46) angegebene formale System in dieser Hinsicht analysiert (S. 79—81 der cit. Abhandlung).

Herr Leśniewski bezweifelt die Strenge des von mir geführten Widerspruchsfreiheitsbeweises für einen gewissen Teil der Mathematik, und konstruiert zur Begründung seiner Ansicht ein „Gegenbeispiel“: er leitet auf Grund der formalen Vorschriften meiner Arbeit zwei Formeln a und $\sim a$ ($\sim$ ist das Zeichen der Negation) her, d. h. einen Widerspruch. Seine Methode ist dabei die folgende (ich erlaube mir seine Schlussweise etwas, wie ich glaube unwesentlich, zu variieren, da es für die folgenden Erläuterungen praktischer ist): Zwei Operationen $O_m^{(1)}(\cdot)$ und $O_n^{(1)}(\cdot)$ werden aufgestellt, für die stets $O_m^{(1)}(a) \neq O_n^{(1)}(a)$ ist ¹⁾, und irgendeine Konstante C_p . Sodann werden $O_m^{(1)}(\cdot)$, $O_n^{(1)}(\cdot)$, C , im Einklang mit meinen „Abänderungs-

¹⁾ D. h. $O_m^{(1)}(a) \neq O_n^{(1)}(a)$ für alle Formeln a bewiesen. Man kann z. B. $O_m^{(1)}(a) = a$, $O_n^{(1)}(a) = \sim a$ erreichen.

vorschriften“ (loc. cit. S. 8—9) in $\square \bullet (\cdot)$, $(\cdot) \bullet \square$, $\square$ abgeändert ¹⁾. Hierdurch gehen (unter Beachtung der „Klammerkonventionen“, loc. cit. S. 9) $O_m^{(1)}(C_p)$ wie $O_n^{(1)}(C_p)$ beide in $\square \bullet \square$ über, obwohl sie $\neq$ sein sollten! Es ist leicht dies zu einem formellen Widerspruch auszugestalten.

Zu diesem Einwande möchte ich bemerken, dass er m. E. auf einer Verkennung des Begriffes „Zeichen“ beruht: wenn man die Mathematik symbolisch beschreibt, so müssen die verschiedenen Zeichen auf alle Fälle voneinander unterscheidbar sein ²⁾. U. zw. besteht offenbar diese „Unterscheidbarkeit“ nicht nur darin, dass zwei als „verschieden“ auszusprechende Zeichen einzeln voneinander unterscheidbar sind, sie umfasst vielmehr notwendigerweise auch die Eigenschaft, dass jede Kombination (lineare Aufeinanderfolge) von Zeichen auf eindeutige Weise in ihre Bestandteile (rein drucktechnisch!) aufgelöst werden kann. Herr Leśniewski hat die Zeichen $\square \bullet$, $\square$, $\bullet \square$ (die ich vorübergehend mit α , β , γ bezeichnen möchte) wohl im Einklang mit meinen allgemeinen Abänderungsvorschriften gewählt, aber dabei gegen ein elementares Gesetz jeder symbolisierenden „Sprache“ verstossen: die Zeichen α , β , γ sind nicht genügend voneinander unterscheidbar, da $\alpha\beta$ mit $\beta\gamma$ identisch ist, d. h. weil es bei der Zeichenkombination $\square \bullet \square$ unmöglich ist zu sagen, ob es sich um $\alpha\beta$ oder $\beta\gamma$ handelt ³⁾.

Es sei hier also nochmals und in aller Schärfe hervorgehoben: Jedes symbolische System, also auch das meine, muss auf solche Zeichen aufgebaut sein, dass zwei Zeichenkombinationen $\alpha\beta\gamma\dots\varrho$ und $\lambda\mu\nu\dots\xi$ nur dann gleich aussehen können, wenn sie aus gleichviel Zeichen bestehen, und α mit λ , β mit μ , γ mit ν , ..., ϱ mit ξ zusammenfällt.

De dies mit dem Gegenstande meiner Arbeit, dem Widerspruchsfreiheitsbeweise für die Mathematik, nichts zu tun hat, sondern einer früheren und m. E. unmathematischen Stufe des Formalismus angehört, glaubte ich in meiner Arbeit hierauf nicht besonders hinweisen zu müssen. Da aber ein Missverständniss entstanden ist, habe ich die Sache jetzt doch erörtert.

¹⁾ Ich schreibe $\square$, $\bullet$ statt 1, $+$ wie Herr Leśniewski, um keine arithmetische Association aufkommen zu lassen.

²⁾ Vgl. D. Hilbert, Hamb. Abh. I, S. 162—163.

³⁾ Dies ändert natürlich nichts an der Tatsache, dass der mathematische Formalismus als prinzipiell sinnlos einzuschätzen ist.

2. Da der Gegenstand einmal angeschnitten ist, möchte ich noch einiges zur Frage der Bezeichnungen sagen, und gleichzeitig ein tatsächliches Versehen in meiner Arbeit richtigstellen¹⁾.

In einer formalistischen Disziplin, wie es z. B. die für die Zwecke der Grundlagen-Untersuchungen formalisierte Mathematik ist, werden die Aussagen durch Formeln ersetzt, d. h. durch Zeichen-Kombinationen, die nach gewissen (rein kombinatorischen) Regeln auseinander erzeugt werden²⁾.

Es ist für die Brauchbarkeit des Systems unerlässlich, dass man einer fertig hingeschriebenen Formel ihre Entstehungsweise eindeutig ansehen kann, dass insbesondere nicht auf zwei verschiedene Weisen dieselbe Formel entstehen können soll³⁾. Dass diese Tatsache die Grundlage aller weiteren (Widerspruchfreiheits-) Überlegungen ist, habe ich auf S. 7 meiner cit. Arbeit gebührend betont.

Dass das von mir angegebene formale System, ohne den „Abänderungen“ (S. 8—9) dieses Postulat erfüllt, habe ich loc. cit. erwähnt, ohne den Beweis auszuführen. Ich habe den, übrigens ziemlich einfachen, Beweis unterdrückt, weil die ganze Sache vom Gesichtspunkte der Widerspruchfreiheitsfrage und der sachlichen Schwierigkeiten derselben aus betrachtet, eine ganz ausserwesentliche Komplikation ist. In der Tat: würde diese Sache, und überhaupt die Bezeichnungsfrage in irgendeiner Beziehung, wesentliche Schwierigkeiten machen, so wäre es ein Leichtes, sie in trivialer Weise ein für allemal aus der Welt zu schaffen. Es würde genügen, statt die Formeln fertig hinzuschreiben, bei jeder Formel ihre Entstehungsgeschichte (im Sinne von Anm. 2)) ausführlich anzugeben. Es ist leicht hierfür eine passende Kurz-Sprache zu finden: z. B. indem man jede Operation mit n Leerstellen durch ein grosses Rechteck

¹⁾ Und auch einen Druckfehler: auf S. 21, Zeile 2 v. o. soll p statt q stehen. Ferner ist auf S. 8, Zeile 10 v. u., natürlich $n > 1$ nötig.

²⁾ So entsteht z. B. die Formel $(x \rightarrow y) \rightarrow ((\sim x \rightarrow y) \rightarrow y)$ aus den Formeln (Variablen) x, y , indem man auf sie $(\cdot \rightarrow \cdot)$ anwendet: $x \rightarrow y$, dann $\sim$ auf x : $\sim x$ dann $(\cdot \rightarrow \cdot)$ auf $\sim x, y$: $(\sim x \rightarrow y)$, dann $(\cdot \rightarrow \cdot)$ auf dieses und y : $((\sim x \rightarrow y) \rightarrow y)$, und schliesslich wieder $(\cdot \rightarrow \cdot)$ auf $(x \rightarrow y)$ und die letztere Formel:

$$((x \rightarrow y) \rightarrow ((\sim x \rightarrow y) \rightarrow y)).$$

Vgl. S. 7 in meiner cit. Arbeit.

³⁾ Bezeichnet man z. F. in der Arithmetik Summe und Produkt mit $\cdot + \cdot$ und $\cdot \times \cdot$, ohne Klammern zu verwenden, so entspricht $x + y \times z$ in der alten Bezeichnungsweise sowohl $(x + y)z$, als auch $x + yz$. Setzt man $x = 1, y = 1, z = 2$, so entsteht der Widerspruch $4 = 3$!

mit dem Seitenverhältniss $1: n + 1$ darstellt, welches in $n + 1$ quadratische Teile eingeteilt ist, und das Operations-Symbol ins erste Fach schreibt, die n Formeln aber, die eingesetzt werden sollen, in die n übrigen Fächer. Wenn man als konsequenter Formalist vor keiner Schwerfälligkeit und Ungewohntheit der Bezeichnungsweise zurückschreckt, ist dies in der Tat ein Allheilmittel.

Tatsächlich ist aber diese radikal-Massregel garnicht notwendig, denn schon die von mir benützte „ungeänderte“ Bezeichnungsweise ist, wie oben gesagt wurde, eindeutig — sie liesse sich sogar, ohne Verlust dieser Eigenschaft, noch wesentlich vereinfachen. Wollte man auf die „Abänderungen“ verzichten, die ich nur darum einführte, um den Anschluss an die übliche, inkonsequente, historisch bedingte, und sich oft ändernde Terminologie der tatsächlichen Mathematik ¹⁾ stets wieder herstellen zu können, so wäre also alles in Ordnung. Rein sachlich liesse sich dagegen auch nichts einwenden.

Nimmt man aber „Abänderungen“ vor, und die vereinfachende „Klammer-Konvention“ (loc. cit. S. 8—9), so muss man sich erneut von der Gültigkeit des Eindeutigkeitssatzes überzeugen. An dieser Stelle habe ich nun loc. cit. tatsächlich zuviel zugelassen. So kann man z. B. drei Operationen $O_r^{(2)}(\cdot, \cdot)$, $O_s^{(1)}(\cdot)$, $O_t^{(1)}(\cdot)$ in $(\cdot \bullet \cdot)$, $\bullet(\cdot)$, $(\cdot) \bullet$ abändern, und dann (mit irgendeiner Konstanten C_p) die Formeln $O_r^{(2)}(C_p, O_s^{(1)}(C_p))$ und $O_r^{(2)}(O_t^{(1)}(C_p), C_p)$ bilden: beide nehmen die Gestalt $(C_p \bullet \bullet C_p)$ an. Um solche Unannehmlichkeiten zu vermeiden, genügt es z. B. den folgenden Zusatz zu meinen „Abänderungs-Vorschriften“ zu machen: kein I' darf gleichzeitig bei einer Abänderung $I'(\cdot, \dots, \cdot)$ und bei einer Abänderung $(\cdot, \dots, \cdot)I'$ Verwendung finden (der zweite Fall könnte sogar auf das Auftreten von nur einer Leerstelle beschränkt werden).

Der Beweis für die Eindeutigkeit abgeänderter Systeme, bei denen diese Vorsichtsmassregel beachtet wird, ist unschwer zu erbringen. Ich glaube aber, dass es sich, mit Rücksicht auf die verschiedenen hier angegebenen Gründe, erübrigt ihn hier durchzuführen — um so mehr, als diese pro-domo-Erörterungen ohnehin schon zu viel Platz einnehmen.

¹⁾ Die 99%, der Mathematiker, ohne Rücksicht auf die Grundlagen-Forschung, de facto verwenden.

Die formalistische Grundlegung der Mathematik

I.

Die Frage nach den Ursachen der allgemein angenommenen unbedingten Zuverlässigkeit der klassischen Mathematik ist durch die kritischen Grundlagen-Untersuchungen der letzten Jahrzehnte, insbesondere aber durch Brouwers System des „Intuitionismus“, erneut aufgerollt worden. Das Bemerkenswerte dabei ist, daß diese, an und für sich philosophisch-erkenntnistheoretische Frage, im Begriffe ist, sich in eine mathematisch-logische zu verwandeln. Die scharfe Formulierung der Mißstände in der klassischen Mathematik durch Brouwer, die exakte und erschöpfende Beschreibung ihrer Methoden (der guten und der bösen) durch Russell, und die von Hilbert geschaffenen Ansätze zur mathematisch-kombinatorischen Untersuchung dieser Methoden und ihrer Zusammenhänge — diese drei wichtigen Vorstöße im Gebiete der mathematischen Logik haben es mit sich gebracht, daß heute in den Grundlagenfragen immer mehr eindeutige mathematische Fragestellungen und nicht Geschmacksunterschiede zu untersuchen sind.

Da die anderen Referate sowohl den durch Brouwer abgegrenzten Bereich der unbedingt und ohne jede Rechtfertigungsnotwendigkeit zuverlässigen, „intuitionistischen“ oder „finiten“ Begriffsbildungen und Beweismethoden, als auch die Russellsche, und in Anlehnung an ihn von seiner Schule weitergeführte, formale Kennzeichnung des Bestandes der klassischen Mathematik erschöpfend behandeln, brauchen wir auf diese nicht näher einzugehen — obwohl ihre Kenntnis selbstverständlich die unerläßliche Voraussetzung zum Verständnis der Zweckmäßigkeit, der Tendenz und des *modus procedendi* der Hilbertschen Beweistheorie ist. Wir wenden uns vielmehr sofort der Beweistheorie zu.

Ihr Grundgedanke ist dieser: Auch wenn die inhaltlichen Aussagen der klassischen Mathematik unzuverlässig sein sollten, so ist

es doch sicher, daß die klassische Mathematik ein in sich geschlossenes, nach feststehenden, allen Mathematikern bekannten Regeln vor sich gehendes Verfahren involviert, dessen Inhalt ist, gewisse, als „richtig“ oder „bewiesen“ bezeichnete, Kombinationen der Grundsymbole sukzessiv aufzubauen. Und zwar ist dieses Aufbauverfahren sicher „finit“ und direkt konstruktiv. Um sich den wesentlichen Unterschied klarzumachen, der zwischen der u. U. unkonstruktiven Behandlung des „Inhalts“ der Mathematik (reelle Zahlen u. ä.), und der immer konstruktiven Verknüpfung ihrer Beweisschritte besteht, vergegenwärtige man sich das folgende Beispiel: Es liege ein klassisch-mathematischer Beweis der Existenz einer reellen Zahl x mit einer gewissen, recht komplizierten und tiefliegenden Eigenschaft $E(x)$ vor. Dann kann es vorkommen, daß man diesem Beweise in keiner Weise ein Verfahren entnehmen kann, ein x mit $E(x)$ zu konstruieren (wir geben gleich ein Beispiel hierfür an); demgegenüber könnte man natürlich, wenn der Beweis die Konventionen des mathematischen Schließens irgendwo verletzen sollte, d. h. wenn in ihm ein Fehler existierte, diesen Fehler durch ein endliches Probiervorgehen sicher finden. D. h.: Die Behauptung eines klassisch-mathematischen Satzes läßt sich nicht immer finit (d. h. überhaupt) kontrollieren, wohl aber der formale Weg, auf dem man zu ihr gelangt. Will man also die klassische Mathematik in bezug auf ihre Zuverlässigkeit prüfen, was prinzipiell nur durch Zurückführung auf das a priori zuverlässige finite System (d. i. das Brouwersche) möglich ist, so darf man nicht ihre Aussagen untersuchen, sondern ihre Beweismethoden. Sie ist als ein kombinatorisches Spiel mit den Grundsymbolen aufzufassen, und es ist kombinatorisch-finit festzustellen, zu welchen Kombinationen der Grundsymbole ihre, „Beweis“ genannten, Aufbaumethoden führen können. —

Wir holen noch das Beispiel eines nichtkonstruktiven Existenzbeweises nach. Sei $f(x)$ eine von 0 bis $\frac{1}{2}$ lineare, von $\frac{1}{2}$ bis $\frac{3}{4}$ lineare, und

von $\frac{3}{4}$ bis 1 lineare Funktion, usw. sei $f(0) = -1$, $f(\frac{1}{2}) = -\sum_{n=1}^{\infty} \frac{\varepsilon_n 2^n}{2^n}$,
 $f(\frac{3}{4}) = \sum_{n=1}^{\infty} \frac{\varepsilon_n 2^n - 1}{2^n}$, $f(1) = 1$. Dabei ist ε_n so definiert: wenn 2^n Summe

von zwei Primzahlen ist, so ist $\varepsilon_n = 0$, sonst $\varepsilon_n = 1$. Offenbar ist $f(x)$ stetig, und an jeder Stelle x beliebig genau effektiv berechenbar. Wegen $f(0) < 0$, $f(1) > 0$ existiert ein x mit $f(x) = 0$ in $0 \leq x \leq 1$

(man erkennt, daß sogar $\frac{1}{2} \leq x \leq \frac{3}{2}$ ist). Trotzdem würde die Angabe einer Wurzel mit einer Genauigkeit $< \frac{1}{2}$ auf große, beim heutigen Stande der Wissenschaft nicht überwundene Schwierigkeiten stoßen: denn sie würde nach sich ziehen, daß man mit Bestimmtheit die Existenz einer Wurzel $< \frac{3}{2}$ bzw. $> \frac{1}{2}$ voraussagen könnte (je nachdem, ob der Näherungswert $\leq \frac{1}{2}$ oder $\geq \frac{1}{2}$ ist). Ersteres schlosse $f(\frac{1}{2}) < 0$, $f(\frac{3}{2}) = 0$ aus, letzteres $f(\frac{1}{2}) = 0$, $f(\frac{3}{2}) > 0$; d. h. ersteres den Fall, daß alle ε_n mit ungeradem n verschwinden, aber nicht alle mit geradem n , letzteres den umgekehrten. D. h. wir hätten bewiesen, daß die bekannte Goldbachsche Vermutung (wonach jedes $2n$ Summe zweier Primzahlen ist), falls sie nicht allgemein zutrifft, schon für ungerade n versagen muß, bzw. schon für gerade n . Und diesen Beweis kann heute noch kein Mathematiker erbringen (weder fürs erste, noch fürs zweite) — als kann auch keiner die Lösung von $f(x) = 0$ genauer als mit dem Fehler $\frac{1}{2}$ bestimmen. (Mit dem Fehler $\frac{1}{2}$ ist $\frac{1}{2}$ ein Näherungswert: denn die Wurzel liegt in $\frac{1}{2}, \frac{3}{2}$, d. h. $\frac{1}{2} - \frac{1}{2}, \frac{1}{2} + \frac{1}{2}$.) —

II.

Die Aufgaben, die die Hilbertsche Beweistheorie zu lösen hat, sind somit die folgenden:

1. Alle Symbole, die in Mathematik und Logik Verwendung finden, aufzuzählen. Sie sollen „Grundsymbole“ heißen. Unter ihnen kommen u. a. die Symbole $\neg$, $\rightarrow$ (die die „Negation“ und die „Implikation“ vertreten) vor.

2. Alle Kombinationen dieser Symbole, welche die in der klassischen Mathematik als „sinnvoll“ bezeichneten Aussagen vertreten, eindeutig zu kennzeichnen. Sie sollen „Formeln“ heißen. (Man beachte: nur „sinnvoll“, nicht etwa auch „richtig“. $1 + 1 = 2$ ist sinnvoll, aber auch $1 + 1 = 1$, ganz unabhängig davon, daß das eine richtig, das andere falsch ist. Sinnlos ist z. B. $1 + \neg = 1$, oder $++1 = -$.)

3. Es ist ein Konstruktionsverfahren anzugeben, das sukzessiv alle Formeln herzustellen gestattet, welche den „beweisbaren“ Behauptungen der klassischen Mathematik entsprechen. Dieses Verfahren heiße darum „Beweisen“.

4. Es ist (finit-kombinatorisch) zu zeigen, daß diejenigen Formeln, welche finit kontrollierbaren (nachrechenbaren arithmetischen) Behauptungen der klassischen Mathematik entsprechen, dann und nur

dann nach 3. bewiesen (d. h. konstruiert) werden können, wenn das soeben erwähnte wirkliche „Nachrechnen“ der ihnen entsprechenden mathematischen Behauptungen die Richtigkeit derselben ergibt.

Sind 1.—4. gesichert, so steht damit die absolute Zuverlässigkeit der klassischen Mathematik für den folgenden Zweck fest: als abkürzende Methode zur Berechnung arithmetischer Ausdrücke, bei welchen das elementare Ausrechnen zu umständlich wäre. Da sie aber immer in diesem Sinne auf die Erfahrung angewandt wird, ist damit die Erfahrungstatsache ihrer Zuverlässigkeit in hinreichendem Maße erklärt.

Nun ist zu beachten, daß 1.—3. nach den Arbeiten Russells und seiner Schule heute gar keine Schwierigkeiten mehr enthalten: Die durch 1.—3. nahegelegte Formalisierung der Mathematik und Logik kann auf viele verschiedene Weisen geleistet werden. Das eigentliche Problem ist 4.

Bei 4. ist nun folgendes zu beachten: Wenn das „effektive Nachrechnen“ einer numerischen Formel deren Richtigkeit ergibt, so läßt sich dies in einen formalen Beweis dieser Formel verwandeln, falls 1.—3. die klassische Mathematik wirklich erschöpfend wiedergeben. Das Kriterium in 4. ist also sicher notwendig, und nur die Hinreichendheit ist zu prüfen. Ergibt nun bei einer numerischen Formel das „effektive Nachrechnen“ deren Unrichtigkeit, so besagt dies, daß man aus ihr eine Beziehung $p = q$ „errechnen“ kann, in der p, q zwei verschiedene, effektiv gegebene Zahlen sind. Wie vorhin, ergäbe dies einen formalen Beweis (nach 3.) für $p = q$. Man erkennt leicht, daß hieraus ein Beweis von $1 = 2$ gewonnen werden kann. Das einzige also, was wir zeigen müssen, um 4. sicherzustellen, ist die formale Unbeweisbarkeit von $1 = 2$ — diese eine spezielle unrichtige numerische Beziehung genügt es zu untersuchen.

Die Unbeweisbarkeit der Formel $1 = 2$ durch die Methoden von 3. wird als „Widerspruchsfreiheit“ bezeichnet. Das eigentliche Problem ist also der finit-kombinatorische Beweis der Widerspruchsfreiheit.

III.

Um die Ansätze zum Widerspruchsfreiheitsbeweise andeuten zu können, müssen wir das Verfahren des formalen Beweisens (nach 3.) etwas näher erläutern. Es wird folgendermaßen definiert:

31. Gewisse Formeln heißen Axiome, sie sind in eindeutiger und finiter Weise gekennzeichnet. Jedes Axiom gilt als bewiesen.

3₂. Sind a , b zwei sinnvolle Formeln, und sind a und $a \rightarrow b$ beide schon bewiesen, so ist auch b bewiesen.

Man beachte: 3₁., 3₂. erlauben wohl, alle beweisbaren Formeln sukzessiv aufzustellen, aber der Prozeß ist nie zu Ende, und sie enthalten kein Verfahren, um von einer gegebenen Formel e zu entscheiden, ob sie beweisbar ist. Denn es ist nicht zu übersehen, welche Formeln sukzessiv bewiesen werden müssen, um schließlich e zu beweisen: es könnten viel kompliziertere und ganz anders gebaute Formeln darunter sein, als e selbst. (Jeder, der z. B. die analytische Zahlentheorie kennt, weiß, wie sehr mit dieser Möglichkeit gerade in den schönsten Kapiteln der Mathematik zu rechnen ist.) Das Problem, durch ein (natürlich finites) allgemeines Verfahren die Beweisbarkeit irgendeiner gegebenen Formel zu entscheiden, ist ein viel tieferes und schwierigeres, als das hier behandelte: es ist das sog. Entscheidungsproblem der Mathematik.

Es würde zu weit führen, die Axiome anzugeben, die in der klassischen Mathematik Verwendung finden. Zur Kennzeichnung genüge folgendes: Es sind unendlich viele Formeln, die als Axiome anzusehen sind (so ist z. B. bei unserer Definition eine jede der Formeln $1 = 1$, $2 = 2$, $3 = 3$, ... Axiom), aber sie entstehen durch Einsetzung aus endlich vielen Schemen. Diese sind so gebaut: „Wenn a , b , c irgendwelche Formeln sind, so ist $(a \rightarrow b) \rightarrow ((b \rightarrow c) \rightarrow (a \rightarrow c))$ Axiom“, u. ä.

Würde es nun gelingen, eine Klasse R von Formeln anzugeben, derart, daß

α) Jedes Axiom zu R gehört,

β) Mit a und $a \rightarrow b$ jedenfalls auch b zu R gehört.

γ) $1 = 2$ nicht zu R gehört,

so wäre die Widerspruchsfreiheit erwiesen: denn offenbar müßte nach α), β) jede bewiesene Formel zu R gehören, nach γ) also $1 = 2$ unbeweisbar sein. Indessen ist an das Angeben einer solchen Klasse R heute nicht zu denken: diese Aufgabe bietet Schwierigkeiten, die mit denjenigen des Entscheidungsproblems vergleichbar sind.

Die folgende Bemerkung leitet indessen von hier zu einer bedeutend einfacheren Aufgabe hinüber: Wenn unser System widerspruchsvoll wäre, so gäbe es einen Beweis von $1 = 2$. Bei diesem Beweise können nur endlich viele Axiome Verwendung gefunden haben, ihre Menge heiße $\mathfrak{M}$. Dann ist also bereits das Axiomensystem $\mathfrak{M}$ widerspruchsvoll. Somit ist das Axiomensystem der klassischen Mathematik gewiß widerspruchsfrei, wenn es jedes endliche

Teilsystem auch ist. Und dies ist nach dem obigen gewiß der Fall, wenn wir zu jeder endlichen Menge $\mathfrak{M}$ von Axiomen eine Formelklasse $R_{\mathfrak{M}}$ angeben können, die die folgenden Eigenschaften besitzt:

α) Jedes Axiom aus $\mathfrak{M}$ gehört zu $R_{\mathfrak{M}}$.

β) Mit a und $a \rightarrow b$ gehört auch b zu $R_{\mathfrak{M}}$.

γ) $1 = 2$ gehört nicht zu $R_{\mathfrak{M}}$.

Hier ist gar keine Verwandtschaft mit dem (viel zu schwierigen) Entscheidungsproblem vorhanden, denn $R_{\mathfrak{M}}$ hängt von $\mathfrak{M}$ ab, und sagt nichts über Beweisbarkeit schlechthin (mit Hilfe aller Axiome) aus.

Es ist selbstverständlich, daß für $R_{\mathfrak{M}}$ eine effektive, finite Konstruktion (für jede effektiv gegebene endliche Axiomenmenge $\mathfrak{M}$) zu fordern ist; und daß die Beweise von α), β), γ) auch finit sein müssen. —

Der gegenwärtige Stand der Dinge ist dadurch gekennzeichnet, daß die Widerspruchsfreiheit der klassischen Mathematik immer noch unbewiesen ist, dagegen dieser Beweis für ein etwas engeres mathematisches System bereits geglückt ist. Dieses ist mit einem System eng verwandt, welches Weyl vor der Aufstellung des intuitionistischen Systems vorgeschlagen hat; dasselbe geht wesentlich über den intuitionistischen Rahmen hinaus, ist aber enger als die klassische Mathematik. (Literaturverweise findet der Leser z. B. in Weyls Artikel über die „Philosophie der Mathematik“ im Handbuch der Philosophie, Oldenbourg, München). Dadurch hat Hilberts System die erste Kraftprobe bestanden: die Rechtfertigung eines nicht finiten und nicht rein konstruktiven mathematischen Systems ist mit finit-konstruktiven Mitteln geglückt. Ob es gelingen wird, diese Rechtfertigung am schwierigeren und wesentlicheren System der klassischen Mathematik zu wiederholen, wird die Zukunft lehren.

Zum Beweise des Minkowskischen Satzes über Linearformen.

Im folgenden soll eine Variante des Beweises von Minkowskis Satz über Linearformen¹⁾ angegeben werden, die bisher wohl noch nicht Beachtung gefunden hat. Ihre Darlegung soll etwa im Anschluß an Minkowskis ursprünglichen Beweis erfolgen, dem der Volumbegriff im n -dimensionalen Raume zugrunde liegt.

Mit Hilfe von Volumbetrachtungen gewinnt bekanntlich Minkowski das folgende Zwischenresultat:

„Es seien n Linearformen

$$L_\mu(x_1, \dots, x_n) = a_{\mu 1} x_1 + \dots + a_{\mu n} x_n \quad (\mu = 1, \dots, n)$$

gegeben, mit der Determinante

$$D = |a_{\mu\nu}| \neq 0.$$

Wenn $c_1, \dots, c_n$ irgendwelche positive Zahlen mit dem Produkt $|D|$ sind, so gibt es mindestens ein System ganzer rationaler $x_1, \dots, x_n$, die nicht alle verschwinden, so daß

$$|L_\mu(x_1, \dots, x_n)| \leq c_\mu \quad (\mu = 1, \dots, n),$$

gilt.“

Dieses Resultat gilt zunächst nur für reelle $a_{\mu\nu}$, wird aber durch eine einfache Transformation auf komplexe übertragen. (In diesem Falle muß mit jedem L_μ auch seine konjugierte Form $\bar{L}_\mu$ vorkommen und $c_\mu = \bar{c}_\mu$ sein.)

Aus Minkowskis Resultat folgt insbesondere

$$|L_1(x_1, \dots, x_n) \cdot \dots \cdot L_n(x_1, \dots, x_n)| \leq |D|.$$

Nun kann für $n > 1$ das Gleichheitszeichen in dieser Relation ausgeschlossen werden (für $n = 1$ ist es, wie man leicht sieht, nicht der Fall), was nach

¹⁾ Minkowski, Gesammelte Abhandlungen, Leipzig 1911, 2. Bd. XXV.

Minkowski durch Grenzbetrachtungen geschieht. Diesen letzteren Schritt auf eine andere, und wie wir glauben einfachere Art auszuführen, ist der Zweck der vorliegenden Note.

Man wähle $c_1 > 0$ von allen Zahlen

$$a_{11}x_1 + \dots + a_{1n}x_n$$

$(x_1, \dots, x_n \text{ ganz rational})$ verschieden; dies ist möglich, da diese Zahlen eine abzählbare Menge bilden. Sodann wähle man $c_2, \dots, c_n$ beliebig, aber unter Berücksichtigung von

$$c_1 \cdot c_2 \cdot \dots \cdot c_n = |D|.$$

Nach dem Zwischenresultat von Minkowski existiert ein System ganzer rationaler $x_1, \dots, x_n$, die nicht alle verschwinden, so daß

$$|L_\mu(x_1, \dots, x_n)| \leq c_\mu \quad (\mu = 1, \dots, n)$$

gilt. Für $\mu = 1$ kann jedoch (nach unserer Annahme über c_1) das Gleichheitszeichen nicht gelten. Folglich ist das Produkt der L_μ kleiner als

$$c_1 \cdot c_2 \cdot \dots \cdot c_n = |D|.$$

ÜBER ADJUNGIERTE FUNKTIONALOPERATOREN

Einleitung.

1. Diese Arbeit ist, wie mehrere frühere¹ den Operatoren des Hilbertschen Raumes gewidmet. Bezeichnungen und Begriffsbildungen lehnen an die in Anm. ¹ genannten Arbeiten des Verfassers an, insbesondere wird die zu verwendende Terminologie in E. (S. 63—78) entwickelt. Demnach ist ein Operator A eine Funktion, die in einer Teilmenge des Hilbertschen Raumes $\mathfrak{H}$ definiert ist und Werte aus $\mathfrak{H}$ annimmt. Weitergehendes wird von Operatoren im allgemeinen nicht verlangt.

Wesentlich sind die folgenden Definitionen (vgl. a. a. O.):

- I. Ein Operator A ist überall sinnvoll, wenn sein Definitionsbereich ganz $\mathfrak{H}$ ist.
- II. A ist Fortsetzung von B , wenn der Definitionsbereich von B Teilmenge des Definitionsbereiches von A ist, und im ersteren A, B übereinstimmen, d. h. $Af = Bf$ gilt. Sind die Definitionsbereiche gleich, so ist $A = B$, sonst ist A echte Fortsetzung von B .
- III. A ist linear, wenn mit $f_1, \dots, f_n$ auch alle $a_1 f_1 + \dots + a_n f_n$ zu seinem Definitionsbereiche gehören² und $A(a_1 f_1 + \dots + a_n f_n) = a_1 A f_1 + \dots + a_n A f_n$ gilt.
- IV. A ist abgeschlossen, wenn für jede Folge $f_1, f_2, \dots$ in seinem Definitionsbereiche, für die f_n gegen ein f und $A f_n$ gegen ein f^* konvergiert ($n \rightarrow \infty$), f auch zum Definitionsbereiche gehört³ und $A f = f^*$ ist.

Die Operatoren, die wir betrachten werden, werden im allgemeinen nicht überall sinnvoll sein, jedoch wird ein gewisser minimaler Umfang des Definitionsbereiches allenfalls erwünscht sein, nämlich dieser:

¹ „Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren“, Math. Annalen 102. 1 (1929), S. 49—131; „Zur Algebra der Funktionaloperationen und Theorie der normalen Funktionaloperatoren“, Math. Annalen 102, 3 (1929), S. 370—427; „Zur Theorie der unbeschränkten Matrizen“, Journal für Math., 161, 4 (1929), S. 208—236; „Über Funktionen von Funktionaloperatoren“, Annals of Math. 32, 2 (1931), S. 191—226. Dieselben werden im Folgenden mehrfach zitiert werden, und zwar abgekürzt als E. bzw. A. bzw. M. bzw. F.

² D. h. der Definitionsbereich ist eine Linear Mannigfaltigkeit.

³ Daß der Definitionsbereich eine abgeschlossene Menge ist, würde hieraus nur folgen, wenn A stetig wäre. Stetigkeit folgte aus der Abgeschlossenheit nur, wenn $\mathfrak{H}$ kompakt wäre, was aber nicht der Fall ist. (Vgl. E., S. 70, insbesondere Anm. ²⁹.)

V. A ist ein $*$ -Operator, wenn die von seinem Definitionsbereich aufgespannte, abgeschlossene Linearmannigfaltigkeit $\mathfrak{S}$ ist⁴, d. h. wenn kein $f \neq 0$ zu seinem Definitionsbereich orthogonal ist.

Wir betrachten nun die Bedingung

$$* \quad (f, Ag) = (f^*, g),$$

wobei f, f^* als fest und g über den ganzen Definitionsbereich von A variabel gedacht sind. Bei gegebenem f mag es u. U. kein $*$ erfüllendes f^* geben, für gewisse f gibt es aber ein f^* : z. B. $f = 0, f^* = 0$. Daß nun $*$ im Lösbarkeitsfalle nur eine Lösung f^* besitze, ist für $*$ -Operatoren charakteristisch⁵, für diese und nur für diese kann daher die folgende Definition aufgestellt werden:

VI. Wenn A ein $*$ -Operator ist, so ist A^* der folgende Operator: Sein Definitionsbereich besteht aus denjenigen f , für die $*$ lösbar ist, u. zw. ist dann $A^*f = f^*$.

Man erkennt sofort, daß A^* , welches in Anlehnung an die übliche Begriffsbildung, die es verallgemeinert⁶, die Adjungierte von A heiße, linear und abgeschlossen ist. Ob es aber außer der 0 noch Stellen gibt, an denen es definiert ist, wissen wir im allgemeinen nicht, noch weniger, ob es ein $*$ -Operator ist. Daß A^{**} (falls es gebildet werden kann, d. h. falls A^* ein $*$ -Operator ist) Fortsetzung von A ist, ist evident.

Diese Begriffsbildungen erlauben, gewisse Begriffe aus E. einfacher zu umschreiben: so ist ein Hermitescher Operator A ein $*$ -Operator, für den A^* Fortsetzung von A ist, ein hypermaximaler Operator einer mit $A^* = A$. (E., S. 70, Def. 6., sowie S. 72, Def. 8., 9.)⁷.

2. Wir werden zeigen, daß die und nur die $*$ -Operatoren A zu linear abgeschlossenen Operatoren B fortgesetzt werden können, für die auch A^* ein $*$ -Operator ist. U. zw. gibt es dann unter allen diesen B einen kleinsten, d. h. einen, von dem alle anderen Fortsetzungen sind, er heiße $\tilde{A}$. Es ist $\tilde{A}^* = A^*$ und $\tilde{A} = A^{**}$. Ein weiteres Kriterium lautet so: Wenn eine Folge von Linearaggregaten $a_1^{(\nu)} f_1^{(\nu)} + \dots + a_n^{(\nu)} f_n^{(\nu)}$ ($\nu = 1, 2, \dots$, $f_1^{(\nu)}, \dots, f_n^{(\nu)}$ aus dem Definitionsbereich) vorliegt derart, daß sie selbst

⁴ Dies wurde in E., S. 70, Def. 6, für Hermitesche und in A., S. 405–406, Def. 5, 6, für normale Operatoren verlangt.

⁵ Denn aus einer Lösung f^* gewinnt man alle anderen $f^* + h$, indem h über alle zum Definitionsbereich von A orthogonalen Elemente von $\mathfrak{S}$ läuft.

⁶ Vgl. E., S. 111–112, insbesondere Anm. ⁶⁴, sowie A., S. 372, Anm. ¹², und S. 405, Def. 5.

⁷ $A^* = A$ steht im Einklang mit der von M. Stone, Proc. Nat. Ac. 16 (1930), S. 173–174, für hypermaximale Operatoren verwendeten Bezeichnung „self-adjoint“.

gegen 0 strebt, und $a_1^{(\nu)} A f_1^{(\nu)} + \dots + a_n^{(\nu)} A f_n^{(\nu)}$ gegen f^* ($\nu \rightarrow \infty$), so ist $f^* = 0$. Wenn A linear ist, lautet dies einfacher: Wenn $f_1, f_2, \dots$ im Definitionsbereich liegen und f_n gegen 0, $A f_n$ gegen f^* strebt ($n \rightarrow \infty$), so ist $f^* = 0^8$.

Diese Operatoren sind von besonderem Interesse:

VII. A ist ein $**$ -Operator, wenn A, A^* beide $*$ -Operatoren sind.

Wenn A ein $**$ -Operator ist, so ist $\tilde{A}$, da es A fortsetzt und $\tilde{A}^* = A^*$ ist, ebenfalls $*$ -Operator — d. h. $\tilde{A}$ ist linear abgeschlossen, $**$ -Operator, Fortsetzung von A und hat dieselbe Adjungierte wie A . Wir gewinnen somit eine vollkommene Übersicht über alle $**$ -Operatoren, wenn wir nur die linear abgeschlossenen unter ihnen betrachten. Diese sind wiederum, nach dem vorhin Gesagten, mit den linear abgeschlossenen $*$ -Operatoren identisch.

Wir werden darum im folgenden stets linear abgeschlossene $*$ -Operatoren betrachten, ist A ein solcher, so ist es nach dem vorhin Gesagten A^* ebenfalls, und es ist $A^{**} = A$ (da $\tilde{A} = A$ ist).

Betrachten wir nun den Operator $A^* A$, d. h. den folgendermaßen definierten: wenn f im Definitionsbereich von A und $A f$ in dem von A^* liegt, ist $A^* A f$ definiert, u. zw. gleich $A^*(A f)$.

Auf Grund der Definition ist klar, daß $A^* A$ Sinn hat und 0 ist, sowie die Linearität von $A^* A$; auch die Abgeschlossenheit ist leicht beweisbar⁹. Wenn $A^* A$ ein $*$ -Operator ist, ist es offenbar Hermitesch und definit¹⁰. Es ist aber vorderhand garnicht zu sehen, ob der Definitionsbereich von $A^* A$ auch andere Elemente hat als 0, oder ob es gar ein $*$ -Operator ist. Klar wäre dies nur, wenn A, A^* überall sinnvoll wären, was aber, wie wir an einen bekannten Satz von Toeplitz anlehnend zeigen werden, nur für stetiges A und A^* der Fall ist.

Es wird uns trotzdem gelingen, allgemein zu zeigen, daß $A^* A$ ein $*$ -Operator, ja sogar hypermaximal ist. Dasselbe gilt, wie Ersetzen von A durch A^* lehrt, für $A A^*$.

Es ist zweckmäßig, diese Verhältnisse durch einige Beispiele zu veranschaulichen. Realisieren wir den Hilbertschen Raum durch alle meß-

⁸ Hieraus folgt die Stetigkeit nicht, weil die Existenz des Limes f^* vorausgesetzt wurde. Vgl. Anm. ³.

⁹ Aus $f_n \rightarrow f$, $A^* A f_n \rightarrow f^*$ folgt: Für $m, n \rightarrow \infty$ ist $A^* A f_m - A^* A f_n \rightarrow 0$, also

$$(A^* A f_m - A^* A f_n, f_m - f_n) = (A f_m - A f_n, A f_m - A f_n) = \|A f_m - A f_n\|^2 \rightarrow 0,$$

also hat $A f_m$ einen Limes, f' . Da A abgeschlossen ist, ist $A f$ sinnvoll und gleich f' , da A^* abgeschlossen ist, ist $A^* f'$ sinnvoll und gleich f^* (man setze $A f_n = f'_n$). Also ist $A^* A f$ sinnvoll und gleich f^* .

¹⁰ $(A^* A f, f) = (A f, A f) = \|A f\|^2$ ist reell und ≥ 0 , vgl. E., S. 70, Def. 6., und S. 74, Def. 10.

baren Funktionen $f(x)$ des Intervalles $0 \leq x \leq 1$ mit endlichem $\int_0^1 |f(x)|^2 dx$ (vgl. z. B. E., S. 108—109); sei $\mathfrak{A}$ die Menge aller stetigen und differenzierbaren $f(x)$ mit endlichem $\int_0^1 |f'(x)|^2 dx$, die gewisse, gleich anzugebende Randbedingungen erfüllen; A in $\mathfrak{A}$ definiert, u. zw. dort $Af = if'$. Als Randbedingungen lassen wir zu:

1. Keine Randbedingungen.
2. $f(1):f(0) = C$ (C eine gegebene Zahl).
3. $f(1) = f(0) = 0$.

Entsprechend sprechen wir von $\mathfrak{A}_1, A_1$; $\mathfrak{A}_{2,C}, A_{2,C}$; $\mathfrak{A}_3, A_3$. $A_1, A_{2,C}, A_3$ sind linear, nicht abgeschlossen, *-Operatoren; A_1 wird von $A_{2,C}$ fortgesetzt, dieses von A_3 , kein $A_{2,C}$ von einem anderen $A_{2,D}$. Aus

$$\begin{aligned}(Af, g) - (f, Ag) &= i \int_0^1 (f'(x) \overline{g(x)} + f(x) \overline{g'(x)}) dx \\ &= i (f(x) \overline{g(x)})_0^1 = i (f(1) \overline{g(1)} - f(0) \overline{g(0)})\end{aligned}$$

folgt, daß A_1^* Fortsetzung von A_3 ist, ebenso $A_{2,C}^*$ von $A_{2,1/\bar{C}}$, und A_3^* von A_1 ; also sind alle drei *-Operatoren, also $A_1, A_{2,C}, A_3$ ***-Operatoren. Also setzen $A_1^* = \tilde{A}_1^*$, $A_{2,C}^* = \tilde{A}_{2,C}^*$, $A_3^* = \tilde{A}_3^*$ bzw. $\tilde{A}_3, \tilde{A}_{2,1/\bar{C}}, \tilde{A}_1$ fort. Man zeigt unschwer, worauf hier nicht näher eingegangen werde, daß diese bzw. Operatoren sogar gleich sind. Daraus folgt: $\tilde{A}_1$ ist Hermiteisch, aber nicht hypermaximal; $\tilde{A}_{2,C}$ ist nur für $C = 1/\bar{C}$, d. h. $|C| = 1$, Hermiteisch, aber dann auch hypermaximal; $\tilde{A}_3$ ist nicht Hermiteisch. Dagegen sind die Operatoren $D_1 = \tilde{A}_1^* \tilde{A}_1 = \tilde{A}_3 \tilde{A}_1$, $D_{2,C} = \tilde{A}_{2,C}^* \tilde{A}_{2,C} = \tilde{A}_{2,1/\bar{C}} \tilde{A}_{2,C}$, $D_3 = \tilde{A}_3^* \tilde{A}_3 = \tilde{A}_1 \tilde{A}_3$ alle drei hypermaximal. Für alle drei ist $Af = -f''$, nur die Definitionsbereiche sind verschieden, die Randbedingungen lauten nämlich:

1. Für D_1 : $f'(1) = f'(0) = 0$.
2. Für $D_{2,C}$: $f(1):f(0) = C$, $f'(1):f'(0) = 1/\bar{C}$.
3. Für D_3 : $f(1) = f(0) = 0$.

Tatsächlich sind dies lauter zulässige Randbedingungen des Eigenwertproblems $-f'' = \lambda f$, was der Hypermaximalität gleichkommt (vgl. E., S. 49, Anm. ², und S. 61).

Allgemeinere bzw. kompliziertere Beispiele (Sturm-Liouvillesche, partielle Differentialgleichungen usw.) zeigen dasselbe Verhalten.

3. Da A^*A und AA^* hermitesch, definit, hypermaximal sind, gibt es (wie sich zeigt) einen und nur einen hermiteschen, definiten, hypermaximalen

Operator B bzw. C mit $B^2 = A^*A$ bzw. $C^2 = AA^*$. Es zeigt sich, daß B denselben Definitionsbereich hat wie A , und stets $\|Bf\| = \|Af\|$ ist; dasselbe ist das Verhältnis von C und A^* . Sodann gelingt die Konstruktion eines Operators W , der im wesentlichen längentreu ist (vgl. w. u.), derart daß

$$A = WB = CW, A^* = BW^* = W^*C, C = WBW^*, B = W^*CW$$

ist. Wie man sieht, ist dies eine Zerlegung von A und A^* , wie sie der Zerlegung einer komplexen Zahl z und ihrer Konjugierten in einen reellen Faktor ≥ 0 (den Absolutwert) und einen Faktor vom Absolutwert 1 entspricht — die hypermaximalen Operatoren B, C benehmen sich ja in vielen Beziehungen wie reelle Zahlen, Definitität entspricht dem ≥ 0 , Längentreue dem Absolutwert 1. Auffallend ist nur, daß zwei verschiedene Operatoren B, C den Absolutwert vertreten, allerdings gehen diese durch die längentreue Abbildung W ineinander über¹¹.

Für überall sinnvolle (also stetige, vgl. w. o.) Operatoren A, A^* ist diese Zerlegung in B, C, W trivial, es ist aber bemerkenswert, daß sie für alle linear abgeschlossenen $*$ -Operatoren A erhalten bleibt und B, C immer hypermaximal ausfallen. Also wird z. B., wenn A hermitesch, definit, aber nicht maximal ist, also A^* echte Fortsetzung von A , also $A^*A = B^2$ Fortsetzung von A^2 , B doch hypermaximal sein — also $B \neq A$, und B nicht etwa Fortsetzung von A , hat es doch denselben Definitionsbereich!

Auf Grund dieser Resultate ist es leicht zu zeigen, daß $B = C$ für die normalen Operatoren A , d. h. diejenigen, die eine komplexe Eigenwertdarstellung zulassen (vgl. A., S. 406, Def. 6., und S. 410—416, Satz 17—19), charakteristisch ist. D. h. $B^2 = C^2$, oder auch $A^*A = AA^*$. Diese Bedingung ist bekanntlich im endlichvieldimensionalen Raume¹² sowie bei vollstetigen Operatoren des Hilbertschen Raumes¹³ für die Normalität charakteristisch, im Hilbertschen Raume zeigte der Verf. noch, daß sie für alle stetigen Operatoren die charakteristische ist (E., S. 115—116, sowie A., S. 406, und der w. o. a. O.), für allgemeine A gewann er aber a. a. O. eine anders gebaute Bedingung. Jetzt ist $A^*A = AA^*$ als allgemein charakteristisch erkannt, ja, es genügt, wenn A^*A Fortsetzung von AA^* ist oder umgekehrt, da beide Operatoren hypermaximal sind.

4. Wenn A irgendein linearer Operator ist, so können wir in seinem Definitionsbereiche $\mathfrak{A}$ durch die Definition $\|f\|_A = \|Af\|$ eine neue Metrik

¹¹ Für vollstetige Operatoren untersuchte E. Schmidt die zu B, C gehörigen Eigenwertprobleme, vgl. Math. Ann. 63 (1907), S. 433—476, §§ 14—16.

¹² Frobenius, Journ. f. Math., 84 (1877), S. 51—54.

¹³ Hellinger und Toeplitz, Enzykl. d. Math. Wiss., II. C. 13., S. 1562—1563.

eingeführen. Damit $\mathfrak{A}$ in $\mathfrak{S}$ überall dicht sei, muß A ein $*$ -Operator sein, damit es im Sinne der neuen Metrik topologisch vollständig sei (wenigstens relativ zu $\mathfrak{S}^{14}$), muß A abgeschlossen sein¹⁵ — d. h. wir werden zu unserer Operatorenklasse (linear, abgeschlossen, $*$) zurückgeführt.

A_1, A_2 bestimmen dieselbe Metrik, wenn sie denselben Definitionsbereich haben, und dort $\|A_1 f\| = \|A_2 f\|$ gilt — wir nennen sie dann metrisch gleich, $A_1 \overline{=} A_2$. Ist der Definitionsbereich von A_1 Teil desjenigen von A_2 und im ersteren $\|A_1 f\| = \|A_2 f\|$, so nennen wir A_2 eine metrische Fortsetzung von A_1 . Sind die Definitionsbereiche gleich, so ist $A_1 \overline{=} A_2$, sonst sei A_2 echte metrische Fortsetzung von A_1 .

Wie wir sahen, ist jedes A einem hermiteschen, definiten, hypermaximalen B metrisch gleich, und man zeigt leicht, daß B hierdurch eindeutig bestimmt ist (vgl. w. u.).

Andererseits werden wir zeigen, daß jeder unstetige Operator A (auch wenn er im gewöhnlichen Sinne unfortsetzbar, d. h. maximal ist), echte metrische Fortsetzungen besitzt. Hier gibt es also keine Maximalität!

5. Unsere Methode wird im wesentlichen darauf beruhen, den Raum aller Paare $\{f, g\}$ (f, g beliebig aus $\mathfrak{S}$) zu betrachten. Wenn wir

$$\begin{aligned} a\{f, g\} &= \{af, ag\}, \\ \{f, g\} + \{h, k\} &= \{f+h, g+k\}, \\ (\{f, g\}, \{h, k\}) &= (f, h) + (g, k) \end{aligned}$$

definieren (also $\|\{f, g\}\|^2 = \|f\|^2 + \|g\|^2$), so ist dies ein Hilbertscher Raum, denn die charakteristischen Bedingungen A.—E. aus E., S. 64—66, gelten für ihn genau so wie für $\mathfrak{S}$. Er verhält sich zu $\mathfrak{S}$ etwa so, wie die Ebene zur Geraden, wir wollen ihn H nennen. Jeder Operator A (in $\mathfrak{S}$) kann in H durch die Menge aller Paare $\{f, Af\}$ dargestellt werden — dies ist das Analogon der graphischen Darstellung der auf der Zahlengeraden definierten Funktionen in der Ebene. Die genannte Menge heiße das Bild von A , B_A .

Man erkennt sofort: $A = B$ bedeutet $B_A = B_B$, A Fortsetzung von B bedeutet B_A umfaßt (als Menge) B_B , A linear bedeutet, daß B_A eine Linearmannigfaltigkeit ist, A abgeschlossen bedeutet, daß B_A (als Menge) abgeschlossen ist.

¹⁴ D. h. daß $\|f_m - f_n\|_A \rightarrow 0$ ($m, n \rightarrow \infty$) die Existenz eines $\|\dots\|_A$ -Limes von $f_1, f_2, \dots$ in $\mathfrak{A}$ nach sich zieht — soweit ein $\|\dots\|$ Limes in $\mathfrak{S}$ existiert. Wenn stets $\|Af\| \geq \varepsilon \cdot \|f\|$ für ein festes $\varepsilon > 0$ gilt, ist die letztere Beschränkung überflüssig.

¹⁵ Diese Betrachtung verschiedener Metrisierungen in $\mathfrak{S}$, so wie der topologischen Vollständigkeit wurde von Herrn K. Friedrichs in einem Vortrage angeregt, den er im Juni 1931 in der Göttinger Math. Ges. hielt.

Wir werden unsere Resultate durch Untersuchungen der B_A in H gewinnen. Die ganze Operatorentheorie ließe sich sogar noch wesentlich konsequenter in H betreiben, als wie es hier geschieht, doch gehen wir hierauf nicht ein.

I. Die Operatoren A , A^* , A^*A .

1. Wir untersuchen, wie in Einl. 4 auseinandergesetzt wurde, den Hilbertschen Raum H und die Bilder B_A in H an Stelle der Operatoren A in §. Wir führen noch den Operator

$$U\{f, g\} = \{ig, -if\}$$

in H ein, er ist Hermitesch und unitär, d. h. $U = U^* = U^{-1}$ (also $U^2 = 1$). Da $(f, Ag) = (f^*, g)$ als $(\{f, f^*\}, U\{g, Ag\}) = 0$ geschrieben werden kann, bedeutet $A^*f = f^*$, daß $\{f, f^*\}$ zu UB_A orthogonal ist, d. h. B_{A^*} ist die Menge aller zu UB_A orthogonalen Elemente von H — d. h. eine abgeschlossene Linearmannigfaltigkeit, also A^* linear abgeschlossen.

Daß es linear abgeschlossene Fortsetzungen von A gibt, bedeutet, daß es abgeschlossene Linearmannigfaltigkeiten B_B gibt, die B_A enthalten. Nun kann jede Menge in H ein B_B sein, wenn sie nie zwei $\{f, f_1^*\}, \{f, f_2^*\}$ mit $f_1^* \neq f_2^*$ enthält — oder, was dasselbe ist, weil es sich um Linearmannigfaltigkeiten handelt, wenn sie nie $\{0, f^*\}$ mit $f^* \neq 0$ enthält. Es ist also bloß zu fordern: die kleinste B_A enthaltende, d. h. die von B_A aufgespannte abgeschlossene Linearmannigfaltigkeit soll kein $\{0, f^*\}$ mit $f^* \neq 0$ enthalten. Ersetzen wir diese Begriffe durch ihre Definitionen, so wird:

SATZ 1. *Der *-Operator A mit dem Definitionsbereich $\mathfrak{A}$ hat dann und nur dann eine linear abgeschlossene Fortsetzung B , wenn Folgendes gilt: Liegen alle $f_1^{(n)}, \dots, f_{\nu_n}^{(n)}$ ($n = 1, 2, \dots$) in $\mathfrak{A}$, und gilt*

$$a_1^{(n)} f_1^{(n)} + \dots + a_{\nu_n}^{(n)} f_{\nu_n}^{(n)} \rightarrow 0, \quad a_1^{(n)} A f_1^{(n)} + \dots + a_{\nu_n}^{(n)} A f_{\nu_n}^{(n)} \rightarrow f^* \quad (n \rightarrow \infty),$$

so ist $f^* = 0$ ¹⁶. U. zw. gibt es dann unter diesen B ein kleinstes, von dem alle anderen Fortsetzungen sind, nämlich dasjenige, für welches B_B die von B_A aufgespannte Linearmannigfaltigkeit ist.

Ein weiteres Kriterium entsteht so. Wenn A^* auch ein *-Operator ist, so existiert auch A^{**} , es ist linear abgeschlossen und offenbar Fortsetzung von A , also existiert $\tilde{A}$. Existiert umgekehrt $\tilde{A}$, so ist $B_{\tilde{A}}$ die von B_A aufgespannte abgeschlossene Linearmannigfaltigkeit, entsprechend $UB_{\tilde{A}}$ und UB_A ,

¹⁶ Die hier stets verwendete starke Konvergenz könnte auch durch schwache ersetzt werden (vgl. z. B. A., S. 378–381), denn für Linearmannigfaltigkeiten (d. i. B_A, B_B) ist die starke Abgeschlossenheit der schwachen, wie E. Schmidt zeigte (Rend. d. Circ. Mat. d. Palermo, 25 (1908), S. 57–73), gleichwertig.

also sind zum ersteren dieselben Elemente von H orthogonal wie zum letzteren — d. h. $\tilde{A}^* = A^*$. Ferner ist $B_{\tilde{A}^*} = B_{A^*}$, die Komplementäre von $UB_{\tilde{A}}$ (es handelt sich jetzt um zwei abgeschlossene Linearmannigfaltigkeiten, vgl. E. S. 74), also, da U unitär, $U^2 = 1$, und die Komplementarität eine symmetrische Beziehung ist, $B_{\tilde{A}}$ die Komplementäre von UB_{A^*} . Das bedeutet in der Terminologie von Einl. 1, daß die dortige Def. VI für A^{**} einen eindeutigen Operator, nämlich $\tilde{A}$ ergibt. Also ist auch A^* ein $*$ -Operator, und $A^{**} = \tilde{A}$. Wir haben bewiesen:

Satz 2. Für die Existenz von $\tilde{A}$ ist auch dieses charakteristisch: A^* ist auch ein $*$ -Operator (für A war das Voraussetzung). U. zw. ist dann $\tilde{A}^* = A^*$, $\tilde{A} = A^{**}$.

Man beachte, daß zwischen den Definitionsbereichen von A und A^* keinerlei Zusammenhänge gefordert oder behauptet wurden.

2. Sei nunmehr A ein linear abgeschlossener $*$ -Operator, dann ist es A^* nach Satz 2 ebenfalls. Da UB_A , B_{A^*} in H komplementär sind, kann jedes Element in H auf eine und nur eine Weise als Summe eines Elementes von UB_A und eines von B_{A^*} dargestellt werden. (vgl. z. B. E., S. 74, Satz 13). Wir wählen das H -Element als $\{0, -if\}$, dann ist

$$\{0, -if\} = \{iAg, -ig\} + \{h, A^*h\}$$

eindeutig lösbar (g im Definitionsbereiche von A , h in dem von A^*). D. h. $iAg + h = 0$, $-ig + A^*h = -if$. Also ist Ag gleich $-ih$, und somit ebenfalls im Definitionsbereiche von A^* , und dann

$$\begin{aligned} -if &= -ig + A^*h = -i(g + A^*Ag) = -i(A^*A + 1)g, \\ (A^*A + 1)g &= f. \end{aligned}$$

Der Operator $A^*A + 1$ ist also für so viele g sinnvoll, daß ihr Wertevorrat der ganze Hilbertsche Raum ist!

Der inverse Operator D , $Df = g$, ist also linear und überall definiert. Ferner ist

$$\begin{aligned} (f, Df) &= (A^*Ag + g, g) = (A^*Ag, g) + (g, g) \\ &= (Ag, Ag) + (g, g) = \|Ag\|^2 + \|g\|^2 \end{aligned}$$

erstens reell, also D Hermitesch. Zweitens ist es ≥ 0 , und nur $= 0$ für $\|g\|^2 = 0$, $g = 0$, $f = (A^*A + 1)g = 0$, also D definit und $Df \neq 0$ für $f \neq 0$. Drittens ist es $\geq \|g\|^2 = \|Df\|^2$, also

$$\|Df\|^2 \leq (f, Df) \leq \|f\| \cdot \|Df\|, \quad \|Df\| \leq \|f\|,$$

also D stetig. Zusammenfassend: D ist ein Hermitescher, definiter, stetiger Operator, aus $f \neq 0$ folgt $Df \neq 0$, das ganze Spektrum von D liegt

zwischen 0 und 1^{17} , und es ist $A^*A + 1 = D^{-1}$. Hieraus ist es ein Leichtes, die Hypermaximalität von $A^*A + 1$ zu folgern, z. B. so. Daß f zu allen Dk orthogonal ist, bedeutet, daß Df zu allen k orthogonal ist (D ist Hermitesch und überall sinnvoll), also $Df = 0$, $f = 0$ — die Komplementäre der vom Wertevorrat von D aufgespannten abgeschlossenen Linearmannigfaltigkeit (in $\mathfrak{H}$) ist (0) , also die letztere $\mathfrak{H}$. Da der Wertevorrat von D der Definitionsbereich von $A^*A + 1$ ist, also auch der von A^*A , ist A^*A ein $*$ -Operator. Nach Einl. 2, insbesondere Anm. ¹⁰, ist also A^*A Hermitesch und definit, daher ist es $A^*A + 1$ auch. Da $A^*A + 1$ alle Werte in $\mathfrak{H}$ annimmt, ist E., S. 102, Satz 41 anwendbar: $A^*A + 1$ ist hypermaximal, also auch A^*A . Also:

Satz 3. *A sei ein linear abgeschlossener $*$ -Operator. Dann ist es A^* auch, und A^*A ist Hermitesch, definit und hypermaximal.*

3. Der Definitionsbereich $\mathfrak{A}_1$ von A^*A ist Teil desjenigen $\mathfrak{A}$ von A , das auf $\mathfrak{A}_1$ eingeschränkte A heiße A_1 — A_1 ist offenbar linear und $*$, aber nicht notwendig abgeschlossen. Da A Fortsetzung von A_1 ist, ist A_1^* Fortsetzung von A^* , also auch $*$, so daß A_1 sogar $**$ ist. Wir wollen $\tilde{A}_1 = A$ beweisen, d. h. daß B_A die von B_{A_1} aufgespannte abgeschlossene Linearmannigfaltigkeit (in H) ist. Anders formuliert: Die Differenz von B_A und der von B_{A_1} aufgespannten abgeschlossenen Linearmannigfaltigkeit ist (0) , oder: in B_A ist nur 0 zur genannten Menge, oder, was dasselbe bedeutet, zu B_{A_1} , orthogonal (vgl. E., S. 74, Def. 11). Ein Element von B_A ist ein (k, Ak) , wenn es zu B_{A_1} orthogonal ist, so ist es insbesondere zu den am Anfang von I., 2. konstruierten $\{g, Ag\}$ orthogonal, also $\{iAk, -ik\}$ zu $\{iAg, -ig\}$ (denn A^*Ag war sinnvoll, als gehört g zu $\mathfrak{A}_1$). Zum dortigen $\{h, A^*h\}$, das zu B_{A^*} gehört, ist das zu UB_A gehörige $\{iAk, -ik\}$ erst recht orthogonal, also auch zu deren Summe, $\{0, -if\}$. Daß $\{iAk, -ik\}$ und $\{0, -if\}$ orthogonal sind (in H), bedeutet, daß es k, f sind (in $\mathfrak{H}$), u. zw. für jedes f — also ist $k = 0$, $\{k, Ak\} = \{0, 0\}$. Damit sind wir am Ziele:

Satz 4. *Der Definitionsbereich von A sei $\mathfrak{A}$, der von A^*A sei $\mathfrak{A}_1$ (Teil von $\mathfrak{A}$), das durch $\mathfrak{A}_1$ eingeschränkte A heiße A_1 . A_1 ist linear, $**$, aber nicht notwendig abgeschlossen. Dagegen ist $\tilde{A}_1 = A$.*

4. Wenn B ein Hermitescher hypermaximaler Operator ist, so gehört nach E., S. 91, Def. 17 und S. 92, Satz 36 eine Zerlegung der Einheit $E(\lambda)$ dazu. Wenn $E(\lambda)$ für $\lambda < 0$ konstant (also $= 0$) ist, so ist nach den dortigen Formeln stets $(Bf, f) \geq 0$, d. h. B definit; ist dagegen $E(\lambda)$ für ein $\lambda < 0$ nicht 0, so gibt es für diese $\lambda = \lambda_0 < 0$ ein $f \neq 0$ mit $E(\lambda_0)f = f$,

¹⁷ Letzteres wegen $0 \leq (f, Df) \leq \|f\|^2$, woraus es durch ähnliche Rechnungen wie A., S. 418, folgt.

also $E(\lambda)f = f$ für $\lambda \geq \lambda_0$ (wegen $E(\lambda_0)$ Teil von $E(\lambda)$ vgl., das in E. über Projektionsoperatoren Gesagte), und dann errechnet man auf Grund derselben Formeln $(Bf, f) \leq \lambda_0 \|f\|^2 < 0$, d. h. B ist nicht definit. Für die Definitität von B ist also $E(\lambda)$ konstant (also $= 0$) für $\lambda < 0$ charakteristisch.

Betrachten wir nun die Schar

$$F(\mu) \begin{cases} = 0 & \text{für } \mu < 0, \\ = E(\sqrt{\mu}) & \text{für } \mu \geq 0. \end{cases}$$

Auch sie ist offenbar eine Zerlegung der Einheit, und der dazugehörige Hermitesche hypermaximale Operator B' ist definit. Wenn $B'f$ Sinn hat, so ist nach Definition

$$\int \mu^2 d\|F(\mu)f\|^2 = \int \mu^2 d\|E(\sqrt{\mu})f\|^2 = \int \lambda^4 d\|E(\lambda)f\|^2$$

endlich, also auch $\int \lambda^2 d\|E(\lambda)f\|^2$ (denn $\int d\|E(\lambda)f\|^2$ ist es ebenfalls und $\lambda^2 \leq \frac{1}{2} + \frac{1}{2}\lambda^4$)—also hat Bf Sinn. Dieselbe Rechnung, wie sie in F., S. 206 ausgeführt wurde (um die dortige Formel e) zu beweisen), zeigt, daß

$$\int \lambda^2 d\|E(\lambda)Bf\|^2 = \int \lambda^4 d\|E(\lambda)f\|^2$$

ist, also endlich—also hat auch $B(Bf)$, d. h. B^2f Sinn. Und wieder eine entsprechende Rechnung zeigt, daß

$$\begin{aligned} (B^2f, g) &= (B(Bf), g) = \int \lambda d(E(\lambda)Bf, g) = \int \lambda^3 d(E(\lambda)f, g) \\ &= \int \mu d(E(\sqrt{\mu})f, g) = \int \mu d(F(\mu)f, g) = (B'f, g) \end{aligned}$$

ist, d. h. $B^2f = B'f$. Daher ist B^2 Fortsetzung des hypermaximalen B' , also $B^2 = B'$.

Etwas anders angeordnet, können wir dies auch so formulieren:

Satz 5. B, B' seien Hermitesche, hypermaximale, definite Operatoren, die bzw. zu ihnen gehörigen Zerlegungen der Einheit $E(\lambda), F(\lambda)$ —also $E(\lambda) = F(\lambda) = 0$ für $\lambda < 0$. Dann ist $B^2 = B'$ mit $E(\lambda) = F(\lambda^2)$ für alle $\lambda \geq 0$ gleichbedeutend.

Wenn also B' gegeben ist und B nicht, so hat $B^2 = B'$ eine einzige derartige Lösung B , sie heiße $\sqrt{B'}$.

Wir wenden nun Satz 5 auf $B' = A^*A$ an und setzen $B = \sqrt{A^*A}$. Der Definitionsbereich von $B^2 = A^*A$ hieß $\mathfrak{A}_1$, das hierauf eingeschränkte A

A_1 , das hierauf eingeschränkte B heiße B_1 (der Definitionsbereich von B umfaßt den von B^* , d. h. $\mathfrak{A}_1$). Nach Satz 3 ist $\tilde{A}_1 = A$, und wenn wir darin A durch B ersetzen (wegen der Hypermaximalität ist $B^* = B$), $\tilde{B}_1 = B$.

A_1, B_1 haben denselben Definitionsbereich, und in ihm ist

$$\|A_1 f\|^2 = (A_1 f, A_1 f) = (A f, A f) = (A^* A f, f)$$

$$\|B_1 f\|^2 = (B_1 f, B_1 f) = (B f, B f) = (B^* B f, f),$$

d. h. $\|A_1 f\| = \|B_1 f\|$. Also sind A_1, B_1 metrisch gleich (vgl. Einl. 4.), und daraus folgt die metrische Gleichheit von $\tilde{A}_1, \tilde{B}_1$ ¹⁸, d. h. von A, B . Einsetzen in die Definition der metrischen Gleichheit ergibt:

Satz 6. Sei A wie in Satz 2., 3., $B = \sqrt{A^* A}$ (vgl. Satz 5.), also Hermitesch, definit, hypermaximal. Dann sind A, B metrisch gleich, d. h.: beide haben denselben Definitionsbereich $\mathfrak{A}$, und in ihm ist stets $\|A f\| = \|B f\|$.

5. Nennen wir den Wertevorrat von A $\mathfrak{B}$, und die abgeschlossenen Hüllen der Linearmannigfaltigkeiten $\mathfrak{A}, \mathfrak{B}$ bzw. $\bar{\mathfrak{A}} = \mathfrak{S}, \bar{\mathfrak{B}}$. Ferner heiße die Menge aller f mit $A f = 0$ $\mathfrak{N}$, sie ist Teil von $\mathfrak{A}$ (und $\bar{\mathfrak{A}}$), und da A linear abgeschlossen ist, eine abgeschlossene Linearmannigfaltigkeit (alles in $\mathfrak{S}$). Die entsprechenden Mengen für A^* seien $\mathfrak{A}', \mathfrak{B}', \mathfrak{N}' = \mathfrak{S}, \bar{\mathfrak{B}}', \mathfrak{N}'$.

f orthogonal zu $\bar{\mathfrak{B}}$ bedeutet dasselbe wie f orthogonal zu $\mathfrak{B}$, d. h. zu allen $A g$, d. h. stets $(f, A g) = 0$ — und dies besagt, daß $A^* f$ definiert ist und $= 0$, d. h. f in $\mathfrak{N}'$. Also: $\mathfrak{N}' = \mathfrak{S} - \bar{\mathfrak{B}}$. Ebenso ist: $\mathfrak{N} = \mathfrak{S} - \bar{\mathfrak{B}}'$. Zusammenfassend:

$$\mathfrak{N} \subset \mathfrak{A}, \mathfrak{N} = \mathfrak{S} - \bar{\mathfrak{B}}', \mathfrak{N}' \subset \mathfrak{A}', \mathfrak{N}' = \mathfrak{S} - \bar{\mathfrak{B}}.$$

Da die Menge der $A f = 0$ nach Satz 4. mit derjenigen der $B f = 0$ identisch ist, ist $\mathfrak{N}$ dieselbe auch für B . Ist der Wertevorrat von B $\mathfrak{B}$ (Linearmannigfaltigkeit!) und seine abgeschlossene Hülle $\bar{\mathfrak{B}}$ (abgeschlossene Linearmannigfaltigkeit!), so ergibt dieselbe Betrachtung $\mathfrak{N} = \mathfrak{S} - \bar{\mathfrak{B}}$ ($B^* = B$ wegen der Hypermaximalität), d. h. $\bar{\mathfrak{B}} = \bar{\mathfrak{B}}'$.

Ersetzen wir A durch A^* , also A^* durch $A^{**} = A$, so entstehen nach Satz 5., 6. die Hermiteschen, definiten, hypermaximalen Operatoren $A A^*$ und $C = \sqrt{A A^*}$. Nach Satz 6. sind A^*, C metrisch gleich, d. h. beide haben denselben Definitionsbereich $\mathfrak{A}'$, und in ihm ist stets $\|A^* f\| = \|C f\|$.

¹⁸ C sei ein linearer $**$ -Operator (A_1, B_1 sind es nach Satz 4.), dann ist B_0 eine Linearmannigfaltigkeit, also B_0 seine abgeschlossene Hülle. D. h. C ist so definiert: Wenn $f_1, f_2, \dots$ im Definitionsbereich von C liegen, $f_n \rightarrow f$ ($n \rightarrow \infty$) und auch die $C f_n$ einen Limes haben — d. h. $\|C(f_m - f_n)\| \rightarrow 0$ ($m, n \rightarrow \infty$) — so liegt f im Definitionsbereich von C , und $C f$ ist der Limes der $C f_n$. Diese Formulierung macht es klar, daß beim Ersetzen von C durch einen metrisch gleichen Operator auch $\tilde{C}$ in einen metrisch gleichen übergeht.

Wir bilden auch für C die Mengen $\mathfrak{B}'$ (Wertevorrat), $\overline{\mathfrak{B}}'$, es ist $\overline{\mathfrak{B}}' = \overline{\mathfrak{B}}$.

Die Zuordnung $Bf \Leftrightarrow Af$ bildet nach Satz 4. $\mathfrak{B}$ eindeutig, linear und längentreu auf $\mathfrak{B}$ ab, sie kann daher eineindeutig, linear und längentreu auf die abgeschlossenen Hüllen $\overline{\mathfrak{B}} = \overline{\mathfrak{B}}'$, $\overline{\mathfrak{B}}$, d. h. $\mathfrak{S} - \mathfrak{N}$, $\mathfrak{S} - \mathfrak{N}'$ ausgedehnt werden. Der Operator dieser längentreuen Abbildung heie V , sein Definitionsbereich ist $\mathfrak{S} - \mathfrak{N}$, sein Wertevorrat $\mathfrak{S} - \mathfrak{N}'$ (vgl. E., S. 70, Def. 6. und S. 71, Satz 10.). Eine ebensolche Abbildung von $\mathfrak{B}'$ auf $\overline{\mathfrak{B}}'$ ist $Cf \Leftrightarrow A^*f$, sie kann ebenso auf $\overline{\mathfrak{B}}' = \overline{\mathfrak{B}}$, $\overline{\mathfrak{B}}'$, d. h. $\mathfrak{S} - \mathfrak{N}'$, $\mathfrak{S} - \mathfrak{N}$ ausgedehnt werden. Ihr längentreuer Operator heie V' , sein Definitionsbereich ist $\mathfrak{S} - \mathfrak{N}'$, sein Wertevorrat $\mathfrak{S} - \mathfrak{N}$. Auf Grund der Definition von V , V' ist

$$A = VB, A^* = V'C.$$

Da es gleichgltig ist, wie V bzw. V' auerhalb des Wertevorrates von B bzw. C , d. h. von $\mathfrak{S} - \mathfrak{N}$ bzw. $\mathfrak{S} - \mathfrak{N}'$ definiert werden, knnen wir sie zu berall definierten linearen Operatoren W bzw. W' fortsetzen, die in $\mathfrak{N}$ bzw. $\mathfrak{N}'$ verschwinden. Dann ist wieder

$$A = WB, A^* = W'C.$$

Wenn wir die Projektionsoperatoren von $\mathfrak{S} - \mathfrak{N}$ bzw. $\mathfrak{S} - \mathfrak{N}'$ E bzw. E' nennen, so ist

$$W = VE, W' = V'E',$$

woraus sofort

$$\|Wf\| = \|Ef\|, \|W'f\| = \|E'f\|$$

folgt, also die Stetigkeit von W , W' . Ferner folgt hieraus

$$\odot \quad W^*W = E, W'^*W' = E'.^{19}$$

Da W' eine Abbildung $\mathfrak{S} - \mathfrak{N} \Leftrightarrow \mathfrak{S} - \mathfrak{N}'$ vermittelt, und $W^*W = E$ die identische Abbildung $\mathfrak{S} - \mathfrak{N} \Leftrightarrow \mathfrak{S} - \mathfrak{N}$, vermittelt W^* die inverse, und ebenfalls längentreue Abbildung $\mathfrak{S} - \mathfrak{N}' \Leftrightarrow \mathfrak{S} - \mathfrak{N}$. Daher ist in $\mathfrak{S} - \mathfrak{N}'$ $\|W^*f\| = \|f\|$. In $\mathfrak{N}'$ ist, da es zum Wertevorrat von W , $\mathfrak{S} - \mathfrak{N}'$, orthogonal ist, $W^*f = 0$ — daher gilt allgemein $W^*E'f = W^*f$, also $\|W^*f\| = \|E'f\|$. Hieraus folgt (vgl. Anm. ¹⁹)

$$\odot \quad WW^* = E',$$

¹⁹ Es ist

$$\|Wf\|^2 = (Wf, Wf) = (W^*Wf, f), \quad \|E'f\|^2 = (E'f, f),$$

also $W^*W = E$; ebenso $W'^*W' = E'$.

ebenso ist

$$\odot \quad W' W'^* = E.$$

Linksmultiplikation von $A = WB$ mit W^* ergibt, da $EB = B$ ist (B hat Werte aus $\mathfrak{S} - \mathfrak{N}$),

$$\ast \quad B = W^* A,$$

ebenso ergibt Linksmultiplikation von $A^* = W' C$ mit W'^*

$$\ast \quad C = W'^* A^*.$$

Wenden wir auf die vier $\ast$ -Formeln $\ast$ an, so zeigt eine einfache Diskussion

$$\begin{array}{ll} \# & A^* = B W^*, \quad A = C W'^*, \\ \# & B = A^* W, \quad C = A W'. \end{array}$$

Bilden wir nun den Operator WBW^{*20} . Da W überall sinnvoll ist, ist er sinnvoll, wenn $W^* f$ in $\mathfrak{A}$ liegt. Da $\mathfrak{N}$ Teil von $\mathfrak{A}$ ist, gehört für jedes Element von $\mathfrak{A}$ seine $\mathfrak{N}$ -Projektion zu $\mathfrak{A}$, also auch seine $\mathfrak{S} - \mathfrak{N}$ -Projektion — d. h. der in $\mathfrak{S} - \mathfrak{N}$ liegende Teil von $\mathfrak{A}$ ist dort überall dicht, d. h. der im Wertevorrat von W^* (dieser ist $\mathfrak{S} - \mathfrak{N}$, vgl. w. o.) liegende Teil von $\mathfrak{A}$ ist dort überall dicht. Da W^* eine Abbildung $\mathfrak{S} - \mathfrak{N}' \Leftrightarrow \mathfrak{S} - \mathfrak{N}$ vermittelt, liegen in $\mathfrak{S} - \mathfrak{N}'$ die f mit $W^* f$ aus $\mathfrak{A}$ überall dicht. Für f in $\mathfrak{N}'$ ist $W^* f = 0$, also auch in $\mathfrak{A}$ — also liegen die genannten f überall dicht, WBW^* ist ein $\ast$ -Operator. Hermitesch und definit ist er ebenso wie B . Auch hypermaximal ist WBW^* , denn sei für alle g , für die es Sinn hat

$$(f, WBW^* g) = (f^*, g),$$

d. h.

$$(W^* f, BW^* g) = (f^*, g).$$

Sei $g = Wh$, h aus $\mathfrak{S} - \mathfrak{N}$, Bh sinnvoll, dann ist $Wg = W^* Wh = h$,

$$(W^* f, Bh) = (f^*, Wh) = (W^* f^*, h).$$

Für jedes h aus $\mathfrak{N}$ hat Bh Sinn, u. zw. 0, und auch rechts steht, da $W^* f^*$ zu $\mathfrak{S} - \mathfrak{N}$ gehört, 0. Also ist stets

$$(W^* f, Bh) = (W^* f^*, h).$$

²⁰ Operatorenprodukte fassen wir konsequent so auf: RSf ist definiert, wenn Sf und $R(Sf)$ definiert sind, und ist gleich $R(Sf)$; entsprechend für QRS , $PQRS$, Diese Multiplikation ist offenbar associativ.

Da B hypermaximal ist, ist BW^*f sinnvoll und gleich W^*f^* , also $WBW^*f = WW^*f^* = f^*$, falls f^* zu $\mathfrak{N}'$ gehört. Dieses folgt aber aus der ursprünglichen Gleichung, da für alle g von $\mathfrak{N}'$ $W^*g = 0$ ist, also deren linke Seite 0 ist.

Nun ist (man beachte die $\times$ -, $\odot$ -, $\#$ -Formeln!)²⁰

$$\begin{aligned}(WBW^*)^2 &= WBW^* \cdot WBW^* = WB \cdot W^*W \cdot BW^* \\ &= WB \cdot E \cdot BW^* = WB \cdot EB \cdot W^* \\ &= WB \cdot BW^* = A \cdot A^* = C^2,\end{aligned}$$

also nach Satz 4 $WBW^* = C$. Nach $\#$ besagt dies $WA^* = C$, und da ein $\times$ $W'^*A^* = C$ aussagt, $WA^* = W'^*A^*$. Somit gilt $Wf = W'^*f$ im ganzen Wertevorrat von A^* , d. i. in $\mathfrak{B}'$ — also auch in $\overline{\mathfrak{B}'} = \mathfrak{S} - \mathfrak{N}$. In $\mathfrak{N}$ gilt es aber auch, da dort beide Seiten gleich 0 sind, also ist $W = W'^*$; Anwenden von $*$ ergibt $W^* = W'$.

Zusammenfassend haben wir also:

SATZ 7. Sei A wie in Satz 2, 3, $B = \sqrt{A^*A}$, $C = \sqrt{AA^*}$ (vgl. Satz 4 mit A und mit A^*), also hermitesch, definit, hypermaximal. A, A^*, B, C mögen die bzw. Linearmanigfaltigkeiten $\mathfrak{A}, \mathfrak{A}', \mathfrak{U}, \mathfrak{U}'$ als Definitionsbereiche, $\mathfrak{B}, \mathfrak{B}', \mathfrak{V}, \mathfrak{V}'$ als Wertevorräte haben, — deute ihre bzw. abgeschlossenen Hüllen an. Die bzw. Mengen (abgeschlossenen Linearmanigfaltigkeiten) $Af = 0$, $A^*f = 0$, $Bf = 0$, $Cf = 0$ seien $\mathfrak{N}, \mathfrak{N}', \mathfrak{R}, \mathfrak{R}'$. Es ist:

$$\begin{aligned}\mathfrak{N} &\subseteq \mathfrak{A} \subseteq \overline{\mathfrak{A}} = \mathfrak{S}, & \mathfrak{N}' &\subseteq \mathfrak{A}' \subseteq \overline{\mathfrak{A}'} = \mathfrak{S}, \\ \overline{\mathfrak{B}'} &= \mathfrak{B} = \mathfrak{S} - \mathfrak{N}, & \overline{\mathfrak{B}} &= \mathfrak{B}' = \mathfrak{S} - \mathfrak{N}'.\end{aligned}$$

Die Projektionsoperatoren von $\mathfrak{S} - \mathfrak{N}$, $\mathfrak{S} - \mathfrak{N}'$ seien bzw. E, E' .

Es gibt einen Operator W (überall definiert, stetig), der $\mathfrak{B}$ auf $\mathfrak{B}'$, also $\mathfrak{S} - \mathfrak{N} = \mathfrak{B}$ auf $\mathfrak{S} - \mathfrak{N}' = \mathfrak{B}'$ eineindeutig, linear, längentreu abbildet; W^* bildet ebenso $\mathfrak{B}'$ auf $\mathfrak{B}$, also $\mathfrak{S} - \mathfrak{N}' = \mathfrak{B}'$ auf $\mathfrak{S} - \mathfrak{N} = \mathfrak{B}$ ab; dabei ist bezüglich der zweitgenannten Abbildungen W^* zu W invers. W ist in $\mathfrak{N}$ Null, W^* ist in $\mathfrak{N}'$ Null. Es ist

$$W^*W = E, \quad WW^* = E'.$$

Es gelten die Beziehungen²⁰:

$$\begin{aligned}A &= WB = CW, & A^* &= BW^* = W^*C, \\ B &= W^*A = A^*W, & C &= AW^* = WA^*,\end{aligned}$$

und aus ihnen folgt

$$C = WBW^*, \quad B = W^*CW.$$

Man beachte: B ist in $\mathfrak{N}$, C in $\mathfrak{N}'$ Null, also nur in $\mathfrak{S} - \mathfrak{N}$ bzw. $\mathfrak{S} - \mathfrak{N}'$ interessant. Hier aber gehen sie durch die längentreue Abbildung von $\mathfrak{S} - \mathfrak{N}$, $\mathfrak{S} - \mathfrak{N}'$ aufeinander, die W , W^* vermitteln, ineinander über. Dies ist die Verallgemeinerung des in Einl. 3, bzw. Anm. 11, erwähnten E. Schmidtschen Satzes.

Wenn 0 weder für A , noch für A^* (Punkt-) Eigenwert ist, d. h. wenn $Af = 0$ ebenso wie $A^*f = 0$ $f = 0$ zur Folge hat, so ist $\mathfrak{N} = \mathfrak{N}' = (0)$, $E = E' = 1$, also $W^*W = WW^* = 1$, d. h. W unitär, $W^* = W^{-1}$. In diesem Falle sind die Formeln von Satz 7 noch durchsichtiger.

II. Anwendungen.

1. Wenn $A^*A = AA^*$ ist, d. h. $B = C$, so folgt aus Satz 7: $A = WB = BW$, $A^* = BW^* = W^*B$, ferner $\mathfrak{N} = \mathfrak{N}'$, $E = F$, also $W^*W = WW^* = E$, also alle Operatoren B , W , W^* miteinander vertauschbar. Hieraus kann die Normalität von A (vgl. Einl. 3) auf verschiedene Weisen gefolgert werden, z. B. könnte man mit den Methoden von E., S. 115 (vor Satz 11*) eine simultane Eigenwertdarstellung für B , W , W^* und damit für A , A^* angeben. Wir schließen uns lieber an Definition in A., S. 116, Def. 6 an und schließen so. A , A^* haben dieselben Definitionsbereiche wie B , C , also beide denselben, sie bilden also im Sinne von A., S. 405, Def. 5 ein konjugiertes Paar. Mit der dortigen Bezeichnungsweise ist, wenn wir noch mit $E(\lambda)$ die zu B gehörige Zerlegung der Einheit bezeichnen,

$$(A)'' = (WB)'' \subseteq (W, B)'' = (W, W^*, \text{alle } E(\lambda))''.$$

Und diese Menge ist Abelsch, weil es die Menge der W , W^* , $E(\lambda)$ ist (vgl. A., S. 389 oben; W , W^* sind mit B vertauschbar, also auch mit allen $E(\lambda)$, A., S. 408, Satz 13).

Ist umgekehrt A normal, so haben die abgeschlossenen Operatoren A , A^* denselben Definitionsbereich (A., S. 416, Satz 21), und wenn A^*Af sowie AA^*g Sinn haben, so ist

$$(f, AA^*g) = (A^*f, A^*g)^{21} = (Af, Ag) = (A^*Af, g).$$

Wegen der Hypermaximalität von AA^* ist also AA^*f sinnvoll und gleich A^*Af . Also ist AA^* Fortsetzung von A^*A , da aber auch dieses hypermaximal ist, $A^*A = AA^*$. Also:

²¹ Nach A., am soeben a. O., ist, wenn Ah Sinn hat, $\|Ah\| = \|A^*h\|$, also $(Ah, Ah) = (A^*h, A^*h)$. Hieraus schließt man, wenn Af , Ag Sinn haben, durch Betrachten von $\frac{f \pm g}{2}$, $\frac{f \pm ig}{2}$, daß auch $(Af, Ag) = (A^*f, A^*g)$ ist.

SATZ 8. Sei A wie in Satz 2, 3, $B = \sqrt{A^*A}$, $C = \sqrt{AA^*}$. Dafür, daß A normal ist, ist $A^*A = AA^*$ charakteristisch oder auch $B = C$. Da A^*A , AA^* beide hypermaximal sind, genügt es sogar, wenn einer den anderen fortsetzt.

Wenn A , A^* dieselben Definitionsbereiche haben und stets $\|Af\| = \|A^*f\|$ ist, so gilt dasselbe für B , C . $Bf \Leftrightarrow Cf$ bildet $\mathfrak{B}$ auf $\mathfrak{B}'$ eineindeutig, linear, längentreu ab, es kann daher ebenso auf $\mathfrak{B} = \mathfrak{S} - \mathfrak{N}$, $\mathfrak{B}' = \mathfrak{S} - \mathfrak{N}'$ ausgedehnt werden; aber es ist $\mathfrak{N} = \mathfrak{N}'$, $\mathfrak{B} = \mathfrak{B}'$. Dieser Operator heiße V_0 , er bildet somit $\mathfrak{B} = \mathfrak{B}' = \mathfrak{S} - \mathfrak{N}$ auf sich ab, wir erweitern ihn zu einem linearen W_0 , das in $\mathfrak{N}$ Null ist. Wie bei Satz 6 ist: $C = V_0 B$, also $C = W_0 B$ und

$$W_0^* W_0 = W_0 W_0^* = E.$$

Anwenden von $*$ ergibt $C = B W_0^*$, also führt W_0^* (innerhalb von $\mathfrak{S} - \mathfrak{N}$) den Definitionsbereich von C in denjenigen von B über, d. h. den von B in sich selbst. Dasselbe gilt für seine Inverse (in $\mathfrak{S} - \mathfrak{N}$) W_0 . Bf und $Cf = W_0 Bf$ (beide sind in $\mathfrak{S} - \mathfrak{N}$) liegen also gleichzeitig im gemeinsamen Definitionsbereich von B, C , d. h. B^2, C^2 haben denselben Definitionsbereich.

Da dort stets

$$(B^2 f, f) = (Bf, Bf) = \|Bf\|^2, \quad (C^2 f, f) = (Cf, Cf) = \|Cf\|^2,$$

also $(B^2 f, f) = (C^2 f, f)$ gilt, ist $B^2 = C^2$ — also $B = C$. Also:

SATZ 9. Für die Normalität von A ist auch dieses charakteristisch: A, A^* haben denselben Definitionsbereich, und in ihm gilt stets $\|Af\| = \|A^*f\|^{22}$.

2. Wir gehen zu den in Einl., 4. definierten Begriffen der metrischen Gleichheit und Fortsetzung von Operatoren über.

Aus Satz 6, sowie der beim Beweise von Satz 9 gemachten Überlegung, die zeigt, daß irgend zwei hermitesche, definite, hypermaximale Operatoren B, C , die metrisch gleich sind, auch wirklich gleich sein müssen (denn der Zusammenhang von B, C mit A, A^* wurde dort nicht benutzt), ergibt sich unmittelbar:

SATZ 10. Jeder linear abgeschlossene $*$ -Operator A ist einem und nur einem hermiteschen, definiten, hypermaximalen B metrisch gleich, d. i. das B von Satz 6: $B = \sqrt{A^*A}$.

Jedes unstetige A hat echte metrische Fortsetzungen. Nach Satz 10 können wir A als hermitesch, definit, hypermaximal annehmen. Nach E., S. 108, Satz 49, ist A echte Fortsetzung eines linear abgeschlossenen

²² In A., S. 423 oben, wurde gezeigt, daß es nicht hinreichend ist, wenn dies bloß auf irgend einer überall dichten Menge gilt.

hermiteschen Operators A' , der sogar auf Kontinuum viele verschiedene Weisen wählbar ist, also ist A'^* (das $*$ ist, da A' $**$ ist) echte Fortsetzung von A . Also:

SATZ 11. *Sei A wie in Satz 10. Wenn er unstetig ist, so hat er echte metrische Fortsetzungen (sogar Kontinuum viele verschiedene).*

Wenn A stetig ist, ist sein Definitionsbereich abgeschlossen, und da er überall dicht sein muß, ist er dann gleich $\mathfrak{D}$. Hieraus, und aus Satz 11 folgt:

SATZ 12. *Sei A wie in Satz 10. A ist dann und nur dann stetig, wenn er überall sinnvoll ist.*

Dies ist die in Einl., 2. erwähnte Verallgemeinerung des Toeplitzschen Satzes²³, die übrigens auch direkt bewiesen werden kann.

Auch über die gleichzeitige metrische Fortsetzung mehrerer Operatoren, z. B. solcher, deren Definitionsbereiche keine gemeinsamen Punkte besitzen (vgl. M., S. 232, Satz 18), gelten sehr allgemeine Sätze, auf die der Verfasser demnächst einzugehen beabsichtigt.

“Proof of the Ergodic Hypothesis”

All footnotes (in the text) 15–28 are incorrect. Their numbers should be increased by one: 16 for 15, 17 for 16, . . . , 29 for 28. The footnotes 1–14 are correct, immediately after footnote 14, footnote 15 should follow.

p. 263, line 5 (from bottom). Formula should read: U_x for U_2 .

p. 265, line 6: read $\mathfrak{H}$ for $\mathfrak{h}$

line 8: read $t - s \rightarrow \infty$ for $t - s \rightarrow 0$

line 18: read $\int_{\Omega} |\mathfrak{F}(f)| dv$ for $\int |\mathfrak{F}(\lambda)|^2 d\lambda$

line 28: read $\mathfrak{N} \subset \mathfrak{M}$ for $\mathfrak{N} \supset \mathfrak{M}$

p. 266, line 6: read c for C

line 14: read $F(\lambda) = 1$ for $F(\lambda) = \Lambda$

line 28: insert footnote 20 at the end of the first sentence.

p. 268, line 9 (formula): read $\int_M$ for $\int_m$

line 20 (formula): read in the suffix $\lambda = \chi_N(P)$ for $\lambda = \chi_n(P)$

line 22 (formula): read $\chi_N^0(P)$ for $X_n^0(P)$

line 2 (from bottom): read $\chi_M \mu M$ for $X_m \mu_m$

line 1 (from bottom) (formula): read $\chi_N^0(P)$ for $X_n^0(P)$

p. 269, line 1: read t_1, s_1 for $t_{1,s1}$

line 3: read $t_{n_v}(P, N) \rightarrow \chi_N^0(P)$ for $t_{n_{qv}}(P, N) \rightarrow X_n^0(P)$

line 4: read $v \rightarrow \infty$ for $v \rightarrow \infty$

line 5: read $\int_M$ for $\int_m$

line 14 (formula): read $\chi_N^0(P) = C_N$ for $X_N^0(P) = C_n$

line 15: read C_v for C_n

line 19 (formula): read C_N for C_n

line 25 (formula): read $\int_{\Omega}$ for $\int$

p. 270, line 1: insert *of finite measure* between *every measurable set* and Λ *remaining invariant*.

p. 271, line 13 (from bottom) (formula): read μ for m , in the numerator and in the denominator.

line 5 (from bottom): read . . . orders . . . for . . . order.

p. 272, line 8 (formula): read $\sum_{v=1}^{k_n}$ for $\sum_{v=1}^{Kn}$ in the denominator.

line 21: read $Z_v^{(n)}$ for $Z_v^{(u)}$

line 22: read $\sum_v^{(m)}$ for $\sum_1^{(n)}$

p. 273, line 10 (footnote): read $(E(\gamma)f, f)$ for $(E(\gamma), f)$

line 21 (footnote): read χ_{Λ} for X_{Λ}

line 22 (footnote): read χ for X or x (in 12 places).

PROOF OF THE QUASI-ERGODIC HYPOTHESIS

1. The purpose of this note is to prove and to generalize the quasi-ergodic hypothesis of classical Hamiltonian dynamics¹ (or "ergodic hypothesis," as we shall say for brevity) with the aid of the reduction, recently discovered by Koopman,² of Hamiltonian systems to Hilbert space, and with the use of certain methods of ours closely connected with recent investigations of our own of the algebra of linear transformations in this space.³ A precise statement of our results appears on page 79.

We shall employ the notation of Koopman's paper, with which we assume the reader to be familiar. The Hamiltonian system of k degrees of freedom corresponding with the Hamiltonian function $H(q_1, \dots, q_k, p_1, \dots, p_k)$ defines a steady incompressible flow $P \longrightarrow P_t = S_t P$ in the space Φ of the variables $(q_1, \dots, q_k, p_1, \dots, p_k)$ or "phase-space," and a corresponding steady conservative flow of positive density ρ in any invariant sub-space $\Omega \subset \Phi$ (Ω being, e.g., the set of points in Φ of equal energy). The Hilbert space $\mathfrak{H}$ consists of the class of measurable functions $f(P)$ having the finite Lebesgue integral $\int_{\Omega} |f|^2 \rho d\omega$, the "inner product"⁴ of any two of them (f, g) and "length" $\|f\|$ being defined by the equations

$$(f, g) = \int_{\Omega} f \bar{g} \rho d\omega; \quad \|f\| = \sqrt{(f, f)}. \quad (1)$$

The transformation U_t is defined as follows:

$$U_t f(P) = f(S_t P) = f(P_t); \quad (2)$$

obviously it has the group property

$$U_t U_s = U_{t+s}, \quad U_0 = I; \quad (3)$$

and in virtue of the conservative character of the flow, and the resulting invariance of $(U_t f, U_t g)$, it is unitary. The spectral reduction of U_t in terms of its "canonical resolution of the identity" $E(\lambda)^5$ is furnished by a theorem due to Stone,⁶ and gives us

$$U_t = \int_{-\infty}^{+\infty} e^{it\lambda} dE(\lambda), \quad (4)$$

this being the symbolic expression for the fact that, for all f, g of $\mathfrak{S}$, we have, in terms of Stieltjès integrals,

$$(U_t f, g) = \int_{-\infty}^{+\infty} e^{it\lambda} d(E(\lambda)f, g). \quad (4')$$

The pith of the idea in Koopman's method resides in the conception of the spectrum $E(\lambda)$ reflecting, in its structure, the properties of the dynamical system—more precisely, those properties of the system which are true "almost everywhere," in the sense of Lebesgue sets.

The possibility of applying Koopman's work to the proof of theorems like the ergodic theorem was suggested to me in a conversation with that author in the spring of 1930. In a conversation with A. Weil in the summer of 1931, a similar application was suggested, and I take this opportunity of thanking both mathematicians for the incentive which they furnished me for undertaking the investigations of this paper.

2: For the sake of brevity, we shall introduce the following notation:

We shall replace $\rho d\omega$ by dv , writing $\int_{\Omega} \rho d\omega = \int_{\Omega} dv$. By the "weight $\mu\theta$ of the Lebesgue-measurable set $\theta (\subset \Omega)$ with respect to the density ρ " will be meant the quantity $\mu\theta = \int_{\theta} \rho d\omega = \int_{\theta} dv$. By a "zero set" we shall mean a set of zero weight, and hence, since throughout Ω , $0 < \rho_1 < \rho < \rho_2$, a set of zero Lebesgue measure.

If θ is a set of points P of Ω or Φ , we shall denote its characteristic function by $\chi_{\theta} = \chi_{\theta}(P)$; i.e., $\chi_{\theta}(P) = 1$ or 0 according as θ does or does not contain P . If $f(P)$ is any measurable function, the set of points for which $f(P) > \lambda$, etc., will be denoted as usual by $[f(P) > \lambda]$, etc. We have the identity $[\chi_{[f > \lambda]} = 1] = [f > \lambda]$, etc.

By the strong convergence of a sequence $f_1, f_2, \dots$ in $\mathfrak{S}$ ($f_n \rightarrow f$) will be meant that $\|f_n - f\| \rightarrow 0$ as $n \rightarrow \infty$. By weak convergence ($f_n \rightharpoonup f$) we mean, on the other hand, that for an arbitrarily chosen g of $\mathfrak{S}$, $(f_n, g) \rightarrow (f, g)$ as $n \rightarrow \infty$. It is shown that $f_n \rightarrow f$ implies $f_n \rightharpoonup f$, but not conversely.⁷ In general, expressions depending for their precise meaning on the nature of the convergence considered will be suffixed by the corresponding convergence symbol, thus we shall write

"separable ($\rightarrow$)," "everywhere dense ($\rightarrow$)," etc. All these notions subsist if n is replaced by one or more continuously varying parameters.

By μ -convergence of a sequence of point sets $\Theta_1, \Theta_2, \dots (\subset \Omega \text{ or } \Phi)$ will be meant the strong convergence of the corresponding measure functions: $\Theta_n \rightarrow \Theta$ if $\chi_{\Theta_n} \rightarrow \chi_{\Theta}$, or, what is the same thing, $\mu[\Theta_n + \Theta - \Theta_n \Theta] \rightarrow 0$ as $n \rightarrow \infty$. Clearly Θ , $\lim \Theta_n$, and $\lim \Theta_n$, will all differ by at most zero sets.

The greatest lower bound and least upper bound of a set $[]$ will be denoted, as usual, by $\inf []$ and $\sup []$.

3. The starting point of our investigations is the construction of the operator

$$\sigma_{t,s} = \frac{1}{t-s} \int_s^t U_{\tau} d\tau \quad (s < t), \quad (5)$$

this being, as before, but the symbolic expression of the fact that for all f, g in $\mathfrak{H}$,

$$(\sigma_{t,s} f, g) = \frac{1}{t-s} \int_s^t (U_{\tau} f, g) d\tau; \quad (5')$$

the existence of $\sigma_{t,s}$ is easily proved.⁸ We will show that, for each f of $\mathfrak{H}$, $\sigma_{t,s} f$ is convergent ($\rightarrow$) as $t-s \rightarrow \infty$, irrespectively of the mode of variation of s, t .

We have from (5'):

$$\begin{aligned} \|\sigma_{t,s} f\|^2 &= (\sigma_{t,s} f, \sigma_{t,s} f) = \frac{1}{t-s} \int_s^t (U_{\tau} f, \sigma_{t,s} f) d\tau \\ &= \frac{1}{(t-s)^2} \int_s^t \int_s^t (U_{\tau} f, U_{\sigma} f) d\sigma d\tau; \end{aligned}$$

since U_{σ} is unitary and $U_{\sigma}^* = U_{\sigma}^{-1} = U_{-\sigma}$,⁹

$$\|\sigma_{t,s} f\|^2 = \frac{1}{(t-s)^2} \int_s^t \int_s^t (U_{\tau-\sigma} f, f) d\tau d\sigma,$$

which reduces, on making the change of variables $\tau - \sigma = x, \tau + \sigma = y$, to

$$\begin{aligned} \frac{1}{(t-s)^2} \int_{-(t-s)}^{+(t-s)} \int_{2s+|x|}^{2t-|x|} (U_x f, f) \frac{dx dy}{2} \\ = \frac{1}{(t-s)^2} \int_{-(t-s)}^{+(t-s)} (t-s-|x|) (U_x f, f) dx. \end{aligned}$$

This may be calculated with the aid of (4'), inasmuch as the various changes in the order of integration are permissible, on account of the uniform convergence of the Stieltjès integral in (4')¹⁰ (for all values of t —at present, x). Thus:

$$\begin{aligned}
\|\sigma_{t,s}f\|^2 &= \frac{1}{(t-s)^2} \int_{-(t-s)}^{+(t-s)} (t-s-|x|) \left[\int_{-\infty}^{+\infty} e^{ix\lambda} d(E(\lambda)f, f) \right] dx \\
&= \frac{1}{(t-s)^2} \int_{-\infty}^{+\infty} \left[\int_{-(t-s)}^{+(t-s)} e^{ix\lambda} (t-s-|x|) dx \right] d(E(\lambda)f, f) \\
&= \frac{2}{(t-s)^2} \int_{-\infty}^{+\infty} \left[\int_0^{t-s} \cos(x\lambda) \cdot (t-s-x) \cdot dx \right] d(E(\lambda)f, f) \\
&= \frac{2}{(t-s)^2} \int_{-\infty}^{+\infty} \frac{1 - \cos(t-s)\lambda}{\lambda^2} d(E(\lambda)f, f) \\
&= \int_{-\infty}^{+\infty} \left[\frac{\sin \frac{1}{2}(t-s)\lambda}{\frac{1}{2}(t-s)\lambda} \right]^2 d(E(\lambda)f, f).
\end{aligned}$$

This integral has a non-negative integrand and a non-decreasing expression after the d -sign; hence we may obtain an upper bound for it as follows: First, break it up into $\int_{-\epsilon}^{+\epsilon}$ and $\int_{-\infty}^{-\epsilon} + \int_{\epsilon}^{+\infty}$; ($\epsilon > 0$, to be considered later). Then replace the integrand in the first part by 1, and that in the second part by $\left[\frac{1}{\frac{1}{2}(t-s)\epsilon} \right]^2 = \frac{4}{(t-s)^2\epsilon^2}$. Finally, replace the field of integration in the second by $\int_{-\infty}^{+\infty}$. We shall then have

$$\begin{aligned}
\|\sigma_{t,s}f\|^2 &\leq \int_{-\epsilon}^{+\epsilon} d(E(\lambda)f, f) + \frac{4}{(t-s)^2\epsilon^2} \int_{-\infty}^{+\infty} d(E(\lambda)f, f) \\
&= \{ (E(\epsilon)f, f) - (E(-\epsilon)f, f) \} + \frac{4}{(t-s)^2\epsilon^2} (f, f).
\end{aligned}$$

Hence, as $t-s \rightarrow +\infty$, $\lim \|\sigma_{t,s}f\|^2 \leq (E(\epsilon)f, f) - (E(-\epsilon)f, f)$, and if, as $\epsilon \rightarrow 0$ $\{E(\epsilon) - E(-\epsilon)\}f \rightarrow 0$, we shall have, on letting $\epsilon \rightarrow 0$, $\|\sigma_{t,s}f\|^2 \rightarrow 0$, so that $\sigma_{t,s}f \rightarrow 0$.

We now introduce the projection operator E_0 defined as follows: $(E(\epsilon) - E(-\epsilon))f \rightarrow E_0f$ as $\epsilon \rightarrow 0$. The existence and projective nature of E_0 is easily deduced from the fact that $E(\epsilon) - E(-\epsilon)$ is a non-increasing "function" of ϵ .¹¹ Thus, we are able to express the condition that $\sigma_{t,s}f \rightarrow 0$ as $t-s \rightarrow 0$ in the form: $E_0f = 0$.

Suppose that $E_0f = f$; then, since $E(\epsilon) - E(-\epsilon) \geq E_0$, we have $E(\epsilon)f = f$, $E(-\epsilon)f = 0$,¹² i.e., $E(\lambda)f = f$ for $\lambda > 0$, $= 0$ for $\lambda < 0$. [Hence, for all g ,

$$(U_t f, g) = \int_{-\infty}^{+\infty} e^{it\lambda} d(E(\lambda)f, g) = (f, g),$$

so that, for all values of t , $U_t f = f$.

Let $\mathfrak{M}$ be the linear manifold in $\mathfrak{S}$ corresponding with E_0 . For every

f of $\mathfrak{M}$, $E_0 f = f$; hence $U_t f = f$, and $\sigma_{t,s} f = f$, so that as $t - s \rightarrow \infty$ we have

$$\sigma_{t,s} f \rightarrow f = E_0 f. \quad (6)$$

For all f orthogonal to $\mathfrak{M}$ on the other hand, we have

$$\sigma_{t,s} f \rightarrow 0 = E_0 f. \quad (6')$$

Now let f be an arbitrary point of $\mathfrak{h}$; we can write $f = f_1 + f_2$, where f_1 is in $\mathfrak{M}$, and f_2 , orthogonal to $\mathfrak{M}$. Then it will follow from (6), (6'), that we still have, as $t - s \rightarrow 0$,

$$\sigma_{t,s} f \rightarrow E_0 f. \quad (6'')$$

Throughout $\mathfrak{M}$, $U_t f = f$. Conversely, if $U_t f = f$, it will follow that $\sigma_{t,s} f = f$, and hence by (6''), $f = E_0 f$, i.e., f belongs to $\mathfrak{M}$. Thus, $\mathfrak{M}$ is the class of all solutions of the equation $U_t f = f$ (i.e., the identity in t).

4. Let us examine $\mathfrak{M}$ more closely. Its elements f are characterized by $U_t f = f$, and hence, in virtue of (2), by

$$f(P) \equiv f(P_t), \quad (7)$$

the $\equiv$ sign holding for all t but with the possible exception of a zero set of points P (in general, dependent on t). Hence, if f is in $\mathfrak{M}$, $\mathfrak{F}(f)$ will be also, provided $\int |\mathfrak{F}(\lambda)|^2 d\lambda$ is finite.

Now f can be expressed as the limit ($\rightarrow$) of functions of the form $\mathfrak{F}(f)$ where $\mathfrak{F}$ is susceptible of but a finite number of values, and these, in their turn, are linear combinations of similar functions susceptible only of the values 0 and 1.¹³ The latter, being of the form $\mathfrak{F}(f)$, belong to $\mathfrak{M}$. If we denote by $\mathfrak{R}$ the class of all functions belonging to $\mathfrak{M}$ and taking on only the values 0 and 1, we may say that $\mathfrak{R}$ spans ($\rightarrow$) the closed ($\rightarrow$) linear manifold $\mathfrak{M}$.

If f belongs to $\mathfrak{R}$, we shall write $[f(P) = 1] = \Lambda$, and $f = f_\Lambda (= \chi_\Lambda)$, (cf. § 2). Since $\mu\Lambda = \|f\|^2$, and f is in $\mathfrak{G}$, $\mu\Lambda$ must be finite; and since f is in $\mathfrak{R} \supset \mathfrak{M}$, it follows from (7) that the transformation $P \rightarrow P_t$ changes Λ by at most a zero set. These two properties, its finite weight and invariance under $P \rightarrow P_t$, characterize Λ . We shall call any set having these properties a Λ -set. Evidently Ω will be a Λ -set if and only if $\mu\Omega$ is finite.

The class of all Λ -sets, being a subclass of the class of all measurable sub-sets of Ω , is separable;¹⁴ that is, there exists a sequence $\Lambda_1, \Lambda_2, \dots$ of Λ -sets such that any Λ -set can be expressed as the limit (in the sense of μ -convergence) of a subsequence of $\Lambda_1, \Lambda_2, \dots$. By methods which we have used in another connection,¹⁵ we are able to replace the sequence $\Lambda_1, \Lambda_2, \dots$ by another sequence of Λ -sets, $\Lambda_1', \Lambda_2', \dots$,

such that, firstly, any member of one sequence may be expressed in terms of elements of the other with the finite repetition of the operations of taking the logical sum (+), the logical product ($\times$) and the logical difference ($-$); secondly, $\Lambda_1', \Lambda_2', \dots$ may be set into one to one correspondence with certain rational numbers $\rho_n = \rho(\Lambda_n')$, $\Lambda_n' = \Lambda(\rho_n)$, $0 < \rho_n < 1$, in such a manner that $\rho_m < \rho_n$ implies $\Lambda(\rho_m) \subset \Lambda(\rho_n)$; and, thirdly, $\inf \rho_n = 0$ and $\sup \rho_n = 1$.

We now define the function $G(P)$ as follows:

$$G(P) \begin{cases} = \inf[\rho_n \text{ for which } P \text{ is in } \Lambda(\rho_n)]; \\ = 1, \text{ when no such } \rho_n \text{ exists.} \end{cases}$$

By its construction, $G(P)$ is invariant under $P \rightarrow P_i$ (remaining unchanged apart from zero sets); and $\inf G(P) = 0$, $\sup G(P) = 1$. Since $[G(P) \leq \rho_n] = \Lambda(\rho_n)$, it follows that $f_{\Lambda_n}' = f_{\Lambda(\rho_n)} = F(G)$ (for we may set $F(\lambda) = \Lambda$ for $\lambda \leq \rho_n$, $= 0$ for $\lambda > \rho_n$); and therefore this property remains true for every f_{Λ_n} ,¹⁶ and any f_Λ of $\mathfrak{R}$ (cf. definition of $\rightarrow$ for sets). And since every f of $\mathfrak{M}$ is the limit ($\rightarrow$) of a sequence of linear combinations of functions of $\mathfrak{R}$, it follows that every such f is a function of G .¹⁷ Finally if $\lambda < \Lambda$, $\mu[G(P) \leq \lambda]$ is finite. For a $\rho_n > \lambda$ may be found, whereupon $[G(P) \leq \lambda] \subset [G(P) \leq \rho_n] = \Lambda(\rho_n) = \Lambda_n'$, and $\mu\Lambda_n' = \|f_{\Lambda_n}'\|^2$, which is finite for any f_{Λ_n}' of $\mathfrak{R}(\subset \mathfrak{S})$.

Any function $G(P)$ like the above, such that $G(P_i) = G(P)$ for all i except perhaps at zero-sets, such that $\lambda' = \inf G(P)$ and $\lambda'' = \sup G(P)$ exist, and that, if $\lambda < \lambda''$, $\mu[G(P) \leq \lambda]$ is finite, and which possesses the property that every f of $\mathfrak{M}$ may be expressed as $\mathfrak{F}(G)$, shall be called a universal integral. We have shown that one universal integral always exists; obviously there are infinitely many.¹⁸

The class $\mathfrak{M}$ coincides with the totality of expressions $\mathfrak{F}(G)$ (with $\int_\Omega |\mathfrak{F}(G(P))|^2 dv$ finite). This makes it possible to express $E_0 f$ in terms of G . We shall give below, instead of our original method of computation, an abbreviated method for which we are grateful to Mr. M. H. Stone.

5. For an arbitrary f of $\mathfrak{S}$, $E_0 f$ belongs to $\mathfrak{M}$, and is, accordingly, of the form $\mathfrak{F}(G)$. Let $\bar{\lambda} < \lambda'' (= \sup G)$, and define $\zeta(\lambda)$ to be 1 for $\lambda < \bar{\lambda}$ and 0 for $\lambda \leq \bar{\lambda}$. Let us set $\Lambda(\lambda) = [G(P) \leq \lambda]$ (we have seen that $\mu\Lambda(\lambda)$ is finite). Then we have, on the one hand,

$$\begin{aligned} \int_\Omega E_0 f(P) \zeta(G(P)) dv &= \int \mathfrak{F}(G(P)) \zeta(G(P)) dv^{20} \\ &= \int_{\lambda'}^{\lambda''} \mathfrak{F}(\lambda) G(\lambda) d\mu\Lambda(\lambda) = \int_{\lambda'}^{\bar{\lambda}} \mathfrak{F}(\lambda) d\mu\Lambda(\lambda), \end{aligned}$$

and, on the other hand,

$$\begin{aligned}
\int_{\Omega} E_0 f(P) \zeta(G(P)) dv &= (E_0 f, \zeta(G)) \\
&= (f, E_0 \zeta(G))^{21} = (f, \zeta(G)) \\
&= \int_{\Omega} f(P) \zeta(G(P)) dv = \int_{\Lambda(\bar{\lambda})} f(P) dv.
\end{aligned}$$

Thus:

$$\int_{\lambda'}^{\bar{\lambda}} \mathfrak{F}(\lambda) d\mu\Lambda(\lambda) = \int_{\Lambda(\bar{\lambda})} f(P) dv. \quad (8)$$

As λ increases from λ' to λ'' , $\mu\Lambda(\lambda)$ goes in a non-decreasing fashion from 0 to $\mu\Omega$; and $\mu\Lambda(\lambda + 0) = \mu\Lambda(\lambda)$. In intervals $\lambda_1 \leq \lambda < \lambda_2$ where $\mu\Lambda(\lambda)$ is constant, $\Lambda(\lambda)$ changes at most by points P of a zero set, i.e., $\lambda_1 < G(P) < \lambda_2$ can be true at most on a zero set; thus the behavior of $\mathfrak{F}(\lambda)$ in such intervals does not affect the relation $E_0 f = \mathfrak{F}(G)$ —we may take $\mathfrak{F}(\lambda)$ constant upon them. It follows that the familiar theorems on the differentiation of Lebesgue integrals may be applied, the independent variable being here $x = \mu\Lambda(\lambda)$. From such considerations it follows that $\frac{d \left\{ \int_{\Lambda(\lambda)} f(P) dv \right\}}{d \left\{ \mu\Lambda(\lambda) \right\}}$ exists for all λ in $\lambda' \leq \lambda < \lambda''$, except for a set of values of λ for which the corresponding set of values $x = \mu\Lambda(\lambda)$ is of zero measure,²² and this derivative is equal to $\mathfrak{F}(\lambda)$. The correspondence $P \rightarrow x$, obtained by setting $\lambda = G(P)$, $x = \mu\Lambda(\lambda)$, carries a set $\Theta \subset \Omega$ into a set θ on the x -axis so that $\mu\theta = \text{measure of } \theta$.²³

Thus we have, except for at most a zero set on Ω ,

$$E_0 f(P) = \left[\frac{d \left\{ \int_{\Lambda(\lambda)} f(P) dv \right\}}{d \left\{ \mu\Lambda(\lambda) \right\}} \right]_{\lambda=G(P)}. \quad (9)$$

This is naturally true for $\lambda = \lambda''$ only when $x = \mu\Lambda(\lambda'')$ exists, i.e., when $\mu\Omega$ is finite.

In the case $\lambda = G(P) = \lambda''$, we carry out the above process with $\zeta(\lambda) = 1$ for $\lambda = \lambda''$, $= 0$ for $\lambda \neq \lambda''$. On setting $\Lambda = [G(P) = \lambda'']$, the following becomes clear: If $\mu\Lambda = 0$, $\mathfrak{F}(\lambda'')$ does not affect $E_0 f(P)$. If $\mu\Lambda = \infty$, $E_0 f(P)$ must be zero; for it is constant on Λ , and belongs to $\mathfrak{S}$. If $\mu\Lambda$ is > 0 and finite, the considerations which lead to (8) show that

$$\mathfrak{F}(\lambda'') \cdot \mu\Lambda = \int_{\Lambda} f(P) dv. \quad (8')$$

Since $\Lambda = \Lambda(\lambda'') - \Lambda(\lambda'' - 0) = \Omega - \Lambda(\lambda'' - 0)$, it follows that

$$E_0 f(P) = \frac{\int_{\Omega - \Lambda(\lambda'' - 0)} f(P) dv}{\mu[\Omega - \Lambda(\lambda'' - 0)]} \text{ (for } G(P) = \lambda''). \quad (9')$$

This formula subsists when $\mu[\Omega - \Lambda(\lambda'' - 0)] = 0$ or ∞ , provided we agree to replace the right-hand member by 0 in the case where the denominator vanishes (and hence the numerator also) or is infinite.

When $\mu\Omega$ is finite, (9') leads to (9) (cf. ²²); otherwise, it forms its natural generalization.

6. Let M and N be any measurable sub sets of Ω , with μM , μN finite. Then $\chi_M(P)$ and $\chi_N(P)$ are in $\mathfrak{S}$, and we may apply our results to them. It follows from (5') that $(\sigma_{t,s} \chi_N, \chi_M)$, which equals $\int_M \sigma_{t,s} \chi_N(P) dv$, is equal to

$$\begin{aligned} & \frac{1}{t-s} \int_s^t \left\{ \int_{\Omega} U_{\tau} \chi_N(P) \cdot \bar{\chi}_M(P) dv \right\} d\tau \\ &= \frac{1}{t-s} \int_s^t \mu(S_{\tau} N \times M) d\tau \quad (S_t P = P_t, \text{ etc.}) \\ &= \int_m Z_{s,t}(P, N) dv, \end{aligned}$$

where we have set $Z_{s,t}(P, N)$ equal to $\frac{1}{t-s}$ times the linear measure of the set of τ -values for which $S_{-\tau} P (= P_{-\tau})$ is on N .²⁴ That is, $Z_{s,t}(P, N)$ is the mean time of sojourn of P in N between the times s and t (actually, with the sign of τ changed, but this is immaterial). Since the above is true for all M , we have

$$\sigma_{t,s} \chi_N(P) = Z_{s,t}(P, N), \quad (10)$$

and in virtue of (6''),

$$Z_{s,t}(P, N) \longrightarrow E_0 \chi_N(P) = \chi_N^0(P) \text{ as } t-s \longrightarrow \infty, \quad (11)$$

in the sense of strong convergence in $\mathfrak{S}$.

On applying (9), (9') to $f = \chi_N$, we have

$$\chi_N^0(P) = \left[\frac{d\{\mu[\Lambda(\lambda) \times N]\}_t}{d\{\mu\Lambda(\lambda)\}} \right]_{\lambda=x_n(P)} \quad (12)$$

when $G(P) < \lambda''$, and

$$X_n^0(P) = \frac{\mu[\Omega - N \times \Lambda(\lambda'' - 0)]}{\mu[\Omega - \Lambda(\lambda'' - 0)]} \quad (12')$$

when $G(P) = \lambda''$. The right-hand members have a meaning except possibly for P on a zero set: in (12), cf.,²² in (12'), we take 0 when the denominator is infinite, or when it (and hence, the numerator) vanishes.

Let us express the content of (11) in the following three ways: (A) explicitly as strong convergence in $\mathfrak{S}$; (B) as point convergence of a subsequence;¹⁷ (C) as weak convergence, or rather as the implication of the latter regarding the inner product of (11) with an arbitrary X_m (μ_m , finite).

$$A. \quad \int [Z_{s,t}(P, N) - X_n^0(P)]^2 dv \longrightarrow 0 \text{ as } t-s \longrightarrow +\infty.$$

B. For every sequence $t_{1,s1}; t_2, s_2; \dots$ with $t_n - s_n \rightarrow +\infty$, there is a sub sequence t_{n_ν}, s_{n_ν} ($\nu = 1, 2, \dots$) such that for all P of Ω with the possible exception of a zero set, $Z_{s_{n_\nu}}, t_{n_{q_\nu}}(P, N) \rightarrow X_N^0(P)$ as $\nu \rightarrow \infty$.

C. For each subset M of Ω of finite μM , $\int_M Z_{s,t}(P, N) dv \rightarrow \int_M X_N^0(P) dv$ as $t - s \rightarrow +\infty$.

We observe that although χ_N^0 is expressed in (12), (12') in terms of the non-uniquely determined universal integral G , its dependence upon G is only apparent: each of (11), A, B or C, determines χ_N^0 uniquely.

7. The existence, for each point P , of the limit of the mean sojourn $Z_{s,t}(P, M)$ is a consequence of A, B or C, and applies to any Hamiltonian system. Our system is *ergodic* if and only if this limit, $\chi_N^0(P)$, is independent of P , i.e., when

$$X_N^0(P) = C_n, \text{ a constant.} \quad (13)$$

When this is true, we must have in the case $\mu\Omega = \infty$ that $C_n = 0$; for otherwise, $\|\chi_N^0\| = \infty$, whereas χ_N^0 is in $\mathfrak{S}$. But when $\mu\Omega$ is finite, we have: $\int \chi_N^0(P) dv = (\chi_N^0, 1) = (E_0 \chi_N, 1) = (\chi_N, E_0 1) = (\chi_N, 1) = \int \chi_N(P) dv = \mu N$. Hence, by (13),

$$C_n = \frac{\mu N}{\mu\Omega}. \quad (13')$$

This is obviously true, from what was said earlier, when $\mu\Omega = \infty$.

It is now a simple matter to tell whether the system is ergodic or not, and we do not even need the more complete results of §§ 4 and 5.

First, suppose that (13) (and consequently (13')) is true for arbitrary N . Let f be an element of $\mathfrak{M}$. Then, on the one hand, we have

$$\int \chi_N^0(P) \bar{f}(P) dv = \frac{\mu N}{\mu\Omega} \int \bar{f}(P) dv = \bar{C}_\mu \cdot N,$$

and, on the other hand,

$$\begin{aligned} \int \chi_N^0(P) \bar{f}(P) dv &= (\chi_N^0, f) = (E_0 \chi_N, f) = (\chi_N, E_0 f) = (\chi_N, f) \\ &= \int \chi_N(P) \bar{f}(P) dv = \int_N \bar{f}(P) dv. \end{aligned}$$

From the equality of the final expressions for all N , we conclude that $f(P) = C$. Secondly, suppose that, conversely, every function of $\mathfrak{M}$ is a constant. Then χ_N^0 , belonging to $\mathfrak{M}$, is a constant, and (13) is true.

Thus the system will be ergodic if and only if $\mathfrak{M}$ consists exclusively of constants. This will be true if and only if $\mathfrak{R}$ consists exclusively of constants,¹³ i.e., that the Λ -sets all differ from 0 or from Ω at most by zero sets. Thus we have proved the theorem:

E. The system is ergodic if and only if every measurable set Λ remaining invariant under S_t (except for points of a zero set) reduces to 0 or to Ω (except for points of a zero set).²⁵

Since in *E* the ergodic condition is the non-existence of any measurable Λ -sets ($\neq 0, \Omega$), one might be tempted to suppose that the ergodic condition as stated in (13) would have to hold for a correspondingly broad class of sets. This, however, is not the case: If (13) is true for all open sets N of finite μN , it will be true, by continuity, for all μ -limits of such sets (cf. §2),—i.e., for all measurable sets N of finite μN .²⁶ Indeed, it is only necessary to require its truth for sets N which are the sums of a finite number of the neighborhoods of an arbitrary "topologically equivalent system of neighborhoods" in Ω ,²⁶ for instance, for sums of finite numbers of spheres.

8. From a purely mathematical standpoint, the question as to the validity and most appropriate generalization of the ergodic hypothesis has been fully answered: these *special* problems have been reduced to the *general* problem of the integrals of the system—the structure of $G(P)$. Thus, the system is either ergodic, or else there is a non-constant $G(P)$, in which case Ω is decomposable into subsets like $[G(P) = \lambda_1]$, $[\lambda_1 \leq G(P) < \lambda_2]$, etc., upon which the flow has a sort of ergodic character, as is easily shown by means of (12), (12').

But from the point of view of physics, there remains the difficult question as to the existence and nature of $G(P)$ in each particular case. It might happen that there are integrals of the system in the classical sense, i.e., analytic, or at least continuously differentiable, as would be true, for example, if $G(P)$ were of this character; in which case they could be used for the reduction of the dimensionality of Ω (cf.,² p. 315, last line). Or it might happen that no such integrals exist, in spite of the fact that $G(P)$ is non-constant. Conceivably this last situation is impossible when the Hamiltonian H is analytic (or even continuously differentiable); if so, the proof of this fact would be most useful. But it appears that the proof could not be obtained alone from the general formal considerations in Koopman's method, i.e., from (3) and from

$$U_t F(f_1(P), f_2(P), \dots) = F(U_t f_1(P), U_t f_2(P), \dots), \quad (14)$$

(cf.,² p. 318). For (4), (14) remain true in the case that the one to one map $P \longrightarrow P_t$ of Ω upon itself is any one-parameter group of the following properties:

- a. P_t is a measurable function of t ,
- b. $P \longrightarrow P_t$ maps every measurable P -set in a measurable P_t -set with the same measure.

Since *a*, *b* permit of discontinuities of P_t , it is easy to give examples with only discontinuous integrals.

Indeed, even when $P \longrightarrow P_l$ is defined by means of equations of the type

$$-\frac{\partial}{\partial t} x_1 = \alpha_1(x_1, \dots, x_l), \dots, \frac{\partial}{\partial t} x_l = \alpha_l(x_1, \dots, x_l), \quad (16)$$

$(P: (x_1, \dots, x_l))$, and when b is valid—when $P \longrightarrow P_l$ is indeed an “incompressible continuous flow”—there are examples where all the integrals are discontinuous, and yet are not constants. (In an example, $l = 2$, α_1/α_2 is continuous in P , but α_1, α_3 themselves, discontinuous.)

We shall not pursue this question further.

We may observe, in conclusion, how remarkable it is that the concept of Lebesgue measure should play so important a rôle in a so essentially physical a question as the validity of the ergodic hypothesis, or, more generally, in the value of the limit of the mean sojourn, $\lim_{t \rightarrow s + \infty} Z(P, N)$.

Even in the case where N is an open set or, indeed, the sum of a finite number of spheres—which has an immediate physical significance—the function $\chi_N^0(P)$ given by the above limit does only need to be measurable! In the last analysis, one is always brought to the cardinal question: “Does P belong to Θ or not?” where the set Θ is merely assumed to be measurable. The opinion is generally prevalent that from the point of view of empiricism such questions are meaningless, for example, when Θ is everywhere dense—for every measurement is of limited accuracy. The author believes, however, that this attitude must be abandoned, and gives the following reason as an argument:

Suppose that Ω , in which P varies, have a finite measure, $m\Omega$. Since Θ is measurable, it follows from a familiar theorem of Lebesgue that

$$\lim_{\epsilon \rightarrow 0} \frac{m[\Theta \times K(P, \epsilon)]}{mK(P, \epsilon)},$$

(where $K(P, \epsilon)$ is a sphere of center P and radius ϵ), exists at each point of Θ , and $= 1$ with the exception of a zero set.²⁷ Similarly for all points of $\Omega - \Theta$, where it is zero, with the same exception. The same is true when the spheres are replaced by many other sorts of figures, e.g., cubes.²⁷ Consider a sequence of partitions of Ω into systems of disjoint cells, $Z_1^{(n)}, \dots, Z_{k_n}^{(n)}$ ($n = 1, 2, \dots$), such that the maximum diameter ϵ_n of $Z_1^{(n)}, \dots, Z_{k_n}^{(n)}$ approaches zero as $n \longrightarrow \infty$. The limited accuracy of measurements finds its expression in the fact that we have to consider different order of accuracy (viz., $1, 2, \dots$); where, by an experiment of order of accuracy n , shall be meant the mere process of distinguishing in which $Z_\nu^{(n)}$ ($\nu = 1, \dots, k_n$) P lies.

Suppose that a measurement of order n has established that, for in-

stance, P lies in $Z_{\nu}^{(n)}$; then the (geometric) probability that P belong to Θ is $\frac{m[Z_{\nu}^{(n)} \times \Theta]}{mZ_{\nu}^{(n)}}$. So that if

$$\frac{m[Z_{\nu}^{(n)} \times \Theta]}{mZ_{\nu}^{(n)}} < \delta \text{ or } > 1 - \delta \quad (\delta > 0), \quad (15)$$

we know with a probability $> 1 - \delta$ the answer to the question, "Is P in Θ or not?" The fact that we will be able to answer this question with a probability $> 1 - \delta$ of being right has the *a priori* probability (i.e., before the observation is made) of

$$w_n^{(\delta)} = \frac{\sum'_{\nu} m Z_{\nu}^{(n)}}{\sum_{\nu=1} m Z_{\nu}^{(n)}} = \frac{\sum'_{\nu} m Z_{\nu}^{(n)}}{m\Omega},$$

where $\sum'_{\nu}$ represents the summation over all values of ν satisfying (15). If we could prove that $w_n^{(\delta)} \rightarrow 1$ as $n \rightarrow \infty$, it would become clear that, granted a sufficiently high accuracy of experiment, the above question could be answered with an arbitrarily great degree of certainty—i.e., the question has physical meaning. (This is seen by taking, e.g., $w_n^{(\delta)} > 1 - \delta$).

Suppose that $w_n^{(\delta)} \rightarrow 1$ as $n \rightarrow \infty$ is untrue. Then for infinitely many values of n , $w_n^{(\delta)} \leq 1 - \eta$ (for a certain $\eta > 0$); so that if $\sum''_{\nu}^{(n)}$ implies summation over all values of ν for which (15) is violated, we shall have

$$m \sum''_{\nu}^{(n)} Z_{\nu}^{(n)} \geq \eta m\Omega.$$

The set Ξ of all points P which belong to infinitely many such sets $\sum''_{\nu}^{(n)} Z_{\nu}^{(n)}$ will then also have a measure $\geq \eta m\Omega > 0$, in virtue of a theorem of Arzela's.²⁸ If P is on Ξ , it lies on infinitely many sets $\sum''_{\nu}^{(n)} Z_{\nu}^{(n)}$; suppose it to be, for example, on $Z_{\nu_n}^{(n)}$. Since ν_n belongs, for infinitely many values of n , to $\sum''_{\nu}^{(n)}$, (15) is violated for these values, the ratio $\frac{m[\Theta \times Z_{\nu_n}^{(n)}]}{mZ_{\nu_n}^{(n)}}$

determines neither the limit 0 nor 1. But this is in contradiction with the theorem of Lebesgue, in the case where the $Z_{\nu}^{(n)}$'s are such that its hypothesis applies (e.g., when $Z_{\nu}^{(n)}$'s are cubes).

¹ For the formulation and critique of this theorem, cf., e.g., *Entykl. d. Math. Wiss.*, 4, Art. 32 on Statistical Mechanics, by P. and T. Ehrenfest, specially 30–36. The original formulations are to be found in *Wien. Ber.*, 63, [2] 679 (1871) (Boltzmann), and *Cambr. Phil. Soc. Trans.*, 12, 547 (1879) (Maxwell).

² These PROCEEDINGS, 17, [5] 315–318 (May, 1931).

³ Cf., e.g., the discussion in the author's paper, "Allgemein Eigenwerttheorie Hermitescher Funktionaloperatoren" (*Math. Ann.*, 102, [1] 108–111 (1929)). This paper, as well as the author's paper, "Zur Algebra der Funktionaloperatoren und Theorie der normalen Operatoren" (*Math. Ann.*, 102, 3 (1929)), will be referred to in the present paper under the abbreviations E and A, respectively.

⁴ Cf. E, 54–55, 109.

⁵ Cf. E, 91–92.

⁶ These PROCEEDINGS, 16, [2] 172–175 (Feb., 1930); also, cf. a paper soon to appear in the *Ann. Math.*

⁷ Cf., regarding these concepts, the article of Hellinger and Toeplitz in the *Math. Encyklopädie*, 2, c. 13, 1435 (1928); further, cf. A, 378–381.

⁸ Cf., e.g., the similar proof in E, 112, top.

⁹ R^* is the adjoint of R in the terminology of matrices: the conjugate-transposed matrix. Cf., e.g., 112.

¹⁰ The integrand, $e^{if\lambda}$, is uniformly bounded, the expression after the d -sign, $(E(\gamma), f)$, of bounded variation: $\int_{-\infty}^{+\infty} d(E(\gamma)f, f) = (f, f)$.

¹¹ E, 91, 77–78.

¹² Cf. the theory of projection operators outlined in E, 74–78; similarly for the discussion to follow.

¹³ Cf. E, 110.

¹⁴ Cf. Hansdorff, *Mengenlehre*, 127 (1927), line 6; or E, 110.

¹⁵ Cf. E, 110.

¹⁶ Cf. the corresponding construction in the proof of theorem 10; in A, 401–402. There, the permutable projections $E_1, E_2, \dots$ and $F_1, F_2, \dots$ took the place of the Λ -sets $\Lambda_1, \Lambda_2, \dots$ and $\Lambda_1', \Lambda_2', \dots$; but this distinction can be abolished by replacing each Λ -set by the operator, $E_\Lambda: f(P) \longrightarrow \chi_\Lambda(P)f(P)$.

¹⁷ For clearly, $X_{\Lambda+M} = \chi_\Lambda + \chi_M - \chi_\Lambda \cdot \chi_M$, $\chi_\Lambda \times M = \chi_\Lambda \cdot \chi_M$, $\chi_{\Lambda-M} = \chi_\Lambda - \chi_\Lambda \chi_M$.

¹⁸ If the sequency $f_1, f_2, \dots$ converges ($\longrightarrow$), a subsequency will converge in any point, excepted a 0-set (cf., f.i. E, 111). Therefore a limit of functions of G is a function of G .

¹⁹ Thus, e.g., each $T(G)$ is one, if $T(\gamma)$ is a monotonically increasing function.

²⁰ On the other hand, our construction shows that it suffices to confine $\mathfrak{F}$ to the second Baire class (the convergence is pointwise convergence except for zero sets).

²¹ This transformation from Lebesgue to Stieltjès integrals goes back to Lebesgue. Cf. *Ann. de l'École Normale*, 3, 27 (1910), p. 407. It is sufficient to establish it for a real variable, for it is then easily extended to an arbitrary Ω , which may always be mapped in a measure-preserving manner upon the real axis. Maps of this sort are given by Lebesgue (loc. cit.) for n -dimensional space, and may easily be extended to Ω .

²² Since $\zeta(G)$ belongs to $\mathfrak{M}$, it is left unchanged by E_0 .

²³ At points of discontinuity of $x = \mu\Lambda(\lambda)$, where x experiences a jump of a whole interval, the differential quotient has a meaning, and is equal to the difference quotient between $\lambda + 0$ and $\lambda - 0$:

$$\frac{\int_{\Lambda(\lambda+0)-\Lambda(\lambda-0)} f(P) dv}{\mu[\Lambda(\lambda+0) - \Lambda(\lambda-0)]}$$

²⁴ Cf. the author's paper, "Über Funktionen von Funktionaloperatoren," *Ann. Math.*, 32, [2] 196 (1931), (Satz 3), as well as the reference (20).

²⁵ This transformation consists in a change in the order of integration; since in the Lebesgue integrals that appear, everything is bounded, it is permissible.

²⁶ For $\mu\Omega = \infty$, naturally only the former will come into question.

²⁷ Cf. E, 110.

²⁸ The generalization to other figures is to be found in its broadest form in Carathéodory, *Vorlesungen über reelle Funktionen* (Leipzig-Berlin, 1918), 492–494, in particular, Theorem 3. The function $f(P)$ appearing there is to be defined as $f(P) = X\Theta(P)$.

²⁹ Cf., e.g., de la Vallée Poussin, *Cours d'Analyse infinitésimale*, 1, 2 (Louvain-Paris, 1909), pp. 68–69.

PHYSICAL APPLICATIONS OF THE ERGODIC HYPOTHESIS

I. In a recent issue of these PROCEEDINGS¹ the author has obtained a proof of the so-called quasi-ergodic hypothesis. The reader is referred to that paper for the precise formulation of this hypothesis, which plays so important a rôle in the foundations of classical statistical mechanics, and thus in the kinetic theory of gasses; the terminology of that paper will be used throughout this note. The exact statement of the mathematical result obtained in the previous paper by the author is as follows:

Let Ω be either the phase-space Φ of the mechanical system considered, or a sub space of Φ invariant under the transformation ($P \longrightarrow P_t$, P a point of Φ , t the time) induced by the equations of motion.² Let dv be the volume element defined in Ω invariant³ under the transformation $P \longrightarrow P_t$, μN the Lebesgue measure (or weight) of $N(\subset \Omega)$ defined by means of $dv : \mu N = \int_N dv$. Let the time of sojourn of P_τ in N during the time $s < \tau < t$, divided by $t - s$, be denoted by $Z_{s,t}(N; P)$.

Then there exists a function $Z(N; P)$ such that, as $t - s \longrightarrow +\infty$, the function of P $Z_{s,t}(N; P)$ converges, in the sense of "strong convergence" in the space of functions of P , to the limit $Z(N; P)$; that is

$$\lim_{t-s \longrightarrow +\infty} \int_{\Omega} |Z_{s,t}(N; P) - Z(N; P)|^2 dv = 0. \quad (1)$$

This property determines $Z(N; P)$, which function, in our previous paper, is studied in more detail and calculated explicitly. The condition for the validity of the so-called quasi-ergodic hypothesis is that $Z(N; P)$ be independent of P ; we have shown in our earlier paper that this will be true if and only if there exists in Ω no integral of the equations of motion

other than one almost everywhere constant (for more details, cf. the paper in question)

This result was obtained on the basis of a method due to B. O. Koopman⁴ in which the study of dynamical systems is undertaken with the aid of certain unitary and Hermitean functional operators, the spectral resolution of which is made use of systematically. Corresponding with the fact that this method operates with reference to function space, is the fact that our result (1) is in terms of "strong convergence," or convergence in the mean. A natural question to ask is whether our results, e.g., (1), could not be established in terms of convergence almost everywhere in Ω .

Mr. G. D. Birkhoff, to whom we communicated these results orally in October, 1931, has subsequently succeeded in establishing the above surmise by means of an extremely astute method of his own in the domain of point set theory; he has proved, i.e., the existence and equality of the numerical limits

$$\lim_{t \rightarrow +\infty} Z_{0,t}(N; P), \quad \lim_{s \rightarrow -\infty} Z_{s,0}(N; P)$$

except on a set of measure zero⁵—which limits, in virtue to (1), will be equal to $Z(N; P)$.

In view of these facts, it is of interest to decide which of the two formulations, (1) or (2), corresponds to the actual physical problem of the ergodic hypothesis. It turns out that the weaker form of statement (1) is sufficient,—that it, indeed, is the precise mathematical equivalent of the physical state of affairs. It is to be noted, further, that the knowledge of the spectral resolution $E(\lambda)$, which is fundamental in Koopman's method,^{1,4} enables one to dominate the physical situation here completely; in particular, it furnishes a numerical estimation of the degree of convergence of the limiting process connected with the ergodic hypothesis, whereas Birkhoff's existence proof for (2) is of a non-constructive character.

II. The physical statement of the problem is as follows:

Consider a function $f(P)$ (in Ω) which is a physical quantity referring to the macroscopic state of the mechanical system (e.g., the pressure of a gas, the mean energy per degree of freedom or temperature, etc.). $f(P)$ changes with the time, having at the instant t the value $f(P_t)$, so that if the interval of time $s < \tau < t$ is so short that the time taken in the measurement fills it completely, the quantity measured is not $f(P)$ itself but its time average for $s < \tau < t$, viz., $\frac{1}{t-s} \int_s^t f(P_\tau) d\tau$. Is it, then, possible to find a constant C (constant with respect to P , C will naturally depend on f) by which $\frac{1}{t-s} \int_s^t f(P_\tau) d\tau$ may be replaced in every application without committing too great an error? (The fact that, if such a C exists

at all, the "micro-canonical" mean $\frac{1}{\mu\Omega} \int_{\Omega} f(P) dv$ may be employed, is evident.)

The criterion for such a possibility is obviously the following: Can a constant C be so determined that the statistical dispersion of $\frac{1}{t-s} \int_s^t f(P_{\tau}) d\tau$ about C is small, i.e.,

$$\int_{\Omega} \left| \frac{1}{t-s} \int_s^t f(P_{\tau}) d\tau - C \right|^2 dv < \epsilon \quad (3)$$

for a given $\epsilon > 0$? Now (1) states precisely that this is the case for $t-s$ sufficiently large, in the case where $f(P)$ is taken to be the characteristic function of the set $N \subset \Omega$:

$$f(P) = \chi_N(P) \begin{cases} = 1, & \text{for } P \text{ in } N \\ = 0, & \text{otherwise.} \end{cases}$$

From this the truth of (3) follows for all $f(P)$, (that is:

$$\lim_{t-s \rightarrow +\infty} \int_{\Omega} \left| \frac{1}{t-s} \int_s^t f(P_{\tau}) d\tau - C \right|^2 dv = 0 \quad (3')$$

(cf. note¹); for the set of functions $f(P)$ for which it is valid forms a closed linear manifold in function space, and such a manifold contains every $f(P)$ if it contains every $\chi_N(P)$ of finite μN (cf. ¹); the integral of $|f(P)|$ must be finite, as is always the case in the applications.

The statement in terms of probability which corresponds to (2) may be made as follows (where the theorem of Egoroff has been used, in virtue of which any almost everywhere convergent sequence of functions will, for an arbitrary $\epsilon > 0$, converge uniformly except for a set of measure $\leq \epsilon^6$):

For every $\epsilon > 0$ and $\delta > 0$ there exists a $T = T(\epsilon, \delta)$ such that the occurrence of $\left| \frac{1}{t} \int_0^t f(P_{\tau}) d\tau - C \right| > \delta$ for any $t > T$ is of probability $\leq \epsilon$.

We have here a refinement of the statement that the statistical dispersion is small; but the latter is quite sufficient for all physical applications.

III. The fact that $t-s$ must be "short" in the physical application and "infinitely long" in the mathematical theorem ($t-s \rightarrow +\infty$ was premised!) can be explained without inconsistency once the degree of convergence in (3') is estimated. Such an estimation may be made with the aid of the method of Koopman.

In the notation of the earlier paper, (3') states that

$$\lim_{t-s \rightarrow +\infty} \| \sigma_{t,s} f - E_0 f \|^2 = 0$$

(cf. note,¹ for this and for the following), and the expression $\| \dots \|^2$ on the left was calculated to be⁷

$$\int_{-\infty}^{-0} + \int_{+0}^{+\infty} \left(\frac{\sin \frac{1}{2} \lambda (t-s)}{\frac{1}{2} \lambda (t-s)} \right)^2 d \| E(\lambda) f \|^2.$$

It is a simple matter to evaluate this expression when $E(\lambda)$ is known. For example, if an interval $-\epsilon < \lambda < \epsilon$ exists, which contains no part of the continuous spectrum, and no point spectrum (besides the simple proper value 0), it is readily seen to be $= \frac{4 \| f \|^2}{\epsilon^2 (t-s)^2}$. In the general case various formulae can be obtained, for example,⁸

$$\begin{aligned} & \left\| \left(E \left(\frac{1}{\sqrt{t-s}} \right) - E(+0) \right) f \right\|^2 \\ & + \left\| \left(E(-0) - E \left(-\frac{1}{\sqrt{t-s}} \right) \right) f \right\|^2 + \frac{4 \| f \|^2}{t-s}, \end{aligned}$$

which furnishes an immediate evaluation. In the case mentioned first $\frac{\| \sigma_{t,s} f - E_0 f \|^2}{\| f \|^2}$ would be $\leq \frac{2}{\epsilon(t-s)}$ this converging uniformly to 0 for all functions f .⁹

These evaluations could easily be analyzed further, but we shall leave the matter now. The example was only given to illustrate the use of Koopman's method in the setting of physical questions.

¹ *Proc. Nat. Acad. Sci.*, **18**, 1 (1932).

² That is, an "integral surface," e.g., an energy surface.

³ In virtue of Liouville's theorem dv is the usual element of volume in the case $\Omega = \Phi$.

⁴ *Proc. Nat. Acad. Sci.*, **17**, 5 (1931).

⁵ *Ibid.*, **17**, 12 (1931).

⁶ *Paris C. R.*, **152** (1911).

⁷ The f appearing there is to be replaced by $f - E_0 f$ (E_0 is permutable with every $E(\lambda)$ and $\sigma_{t,s}$); this has for effect that $\int_{-\infty}^{-0} + \int_{+0}^{+\infty}$ replaces the earlier $\int_{-\infty}^{+\infty}$.

⁸ Separate $\int_{-\infty}^{-0} + \int_{+0}^{+\infty}$ into $\left[\int_{+0}^{+\frac{1}{\sqrt{t-s}}} + \int_{-\frac{1}{\sqrt{t-s}}}^{-0} \right] + \left[\int_{+\frac{1}{\sqrt{t-s}}}^{+\infty} + \int_{-\infty}^{-\frac{1}{\sqrt{t-s}}} \right]$

and not that $\left(\frac{\sin \frac{1}{2} \lambda (t-s)}{\frac{1}{2} \lambda (t-s)} \right)^2 \leq 1$ and $\leq \left(\frac{2}{\lambda (t-s)} \right)^2$.

⁹ It is easily shown, that this only occurs for periodical motions in which all paths have the same period.

DYNAMICAL SYSTEMS OF CONTINUOUS SPECTRA

1. In a recent paper by B. O. Koopman,¹ classical Hamiltonian mechanics is considered in connection with certain self-adjoint and unitary operators in Hilbert space $\mathfrak{H}$ ($= \mathfrak{L}_2$). The corresponding canonical resolution of the identity $E(\lambda)$, or "spectrum of the dynamical system," is introduced, together with the conception of the spectrum revealing in its structure the mechanical properties of the system.² In general, $E(\lambda)$ will consist of a discontinuous part (the "point spectrum") and of a continuous part. The case of a pure point spectrum, and the other extreme, that in which the inner product $(E(\lambda)f, g)$ is, for every f and g in $\mathfrak{H}$, the Lebesgue integral of one of its derivatives, may readily be treated by known analytical tools.³ The present paper is devoted to the case where $E(\lambda)$ is continuous ($\lambda \neq 0$), but without $(E(\lambda)f, g)$ being necessarily equal to the integral of its derivative. It will further be assumed that the system is non-integrable in the sense that any f in $\mathfrak{H}$ such that $U_t f = f$ almost everywhere on Ω must be almost everywhere constant. In other words, we are assuming the following hypothesis:

$$C. \quad E(\lambda + 0) - E(\lambda - 0) = 0, \text{ for } \lambda \neq 0:$$

If $E_0 = E(+0) - E(-0)$, then $E_0 f =$ almost everywhere constant.

Suppose that ϕ is a characteristic function of a non-integrable system: $U_t\phi = e^{i\lambda t}\phi$. Then $U_t|\phi| = |\phi|$, so that $|\phi|$ must be almost everywhere constant, and hence we may take $\phi = e^{i\theta}$, θ being defined almost everywhere (mod 2π).⁴ From $U_t\phi = e^{i\lambda t}\phi$ follows that $U_t\theta = \theta + \lambda t$ (mod 2π); thus θ is an "angle variable." Since ϕ is in $\mathfrak{S}$, $\|\phi\| = \int_{\Omega} e^{i\theta} e^{-i\theta} dv = \mu\Omega$ must be finite, so that a non-integrable system of infinite $\mu\Omega$ can have no characteristic function. These considerations permit the following restatement of our hypothesis:

C'. The system possesses no invariant subset of Ω of positive finite measure, and in the case where $\mu\Omega$ is finite, no angle variables.

We will show that under this hypothesis all the initially observed properties of the system are obliterated by the lapse of time: the method of elementary mechanics of computing the final from the initial state must be replaced by the methods of the theory of probability. Contrary to the case of classical statistical mechanics, this situation is not dependent upon the system's having an enormous number of degrees of freedom: this number may perfectly well reduce to two.

2. By a P -set we shall mean a set of points of Ω , by a t -set, a set of points on the time axis. The P -set for which $f(P) > a$, etc., shall be denoted as usual by $[f(P) > a]$, etc., and similarly for t -sets and functions of t . Along with the μ -measure μM of a P -set M defined with respect to the volume element $dv = \rho d\omega$, we shall define the τ -measure τI of a t -set I as follows:

$$\tau I = \lim_{T \rightarrow +\infty} \frac{m(I \cdot [|t| \leq T])}{2T}.$$

Here m denotes the ordinary Lebesgue measure, and the customary set-multiplication is referred to.⁵ By a zero P -set or a zero t -set we shall mean one for which the μ -measure or the τ -measure, respectively, is zero. The "characteristic function" of a P -set or a t -set Θ shall be denoted, as usual, by χ_{Θ} : $\chi_{\Theta}(P) = 1$ or 0 according as P is or is not on Θ , etc. In these terms we are able to state the fundamental theorem of this paper:

THEOREM I. Under the hypothesis *C*, there exists a zero t -set I such that, for any two P -sets M and N of finite μ -measure, as $t \rightarrow +\infty$ (or $t \rightarrow -\infty$) through values not on I ,

$$\mu(M_t \cdot N) \rightarrow \frac{\mu M \cdot \mu N}{\mu\Omega}, \quad (1)$$

M_t being the image of M after the lapse of time t . When $\mu\Omega = \infty$, the right-hand member is to be replaced by zero.

Proof.—It has been shown⁶ that in non-integrable systems,

$$(E_0 \chi_M, \chi_N) = \frac{\mu M \cdot \mu N}{\mu \Omega};$$

hence (1) is equivalent to

$$\lim_{\substack{t \rightarrow \pm \infty \\ t \text{ not on } I}} (U_t \chi_M, \chi_N) = (E_0 \chi_M, \chi_N),$$

which, in turn, may be replaced by

$$\lim_{\substack{t \rightarrow \pm \infty \\ t \text{ not on } I}} (U_t f, g) = (E_0 f, g), \quad (1')$$

since the aggregate of χ -functions and their linear combinations is everywhere dense in $\mathfrak{H}$. That is, we must prove that as $t \rightarrow \pm \infty$ through values not on I , $U_t f$ converges "weakly" to $E_0 f$.⁷ Now $(U_t f, g) \rightarrow (E_0 f, g)$ for all f, g if and only if $(U_t f, f) \rightarrow (E_0 f, f)$ for all f .⁸ Using, now, the known properties:⁹ $U_t E_0 = E_0 U_t = U_t$, $E_0^2 = E_0$, etc., we have $((U_t - E_0)f, f) = ((U_t - U_t E_0)f, f) = (U_t(I - E_0)f, f) = (U_t(I - E_0)^2 f, f) = (U_t(I - E_0)f, (I - E_0)f)$, so that our problem is to show that $(U_t(I - E_0)f, (I - E_0)f) \rightarrow 0$. That is, we are to show that a zero t -set I exists such that, as $t \rightarrow \pm \infty$, t not on I , $(U_t g, g) \rightarrow 0$ whenever $E_0 g = 0$. Since U_t is a bounded operator, it is sufficient to show this only for an everywhere dense sequence $g_1, g_2, g_3, \dots$ of functions g .

We now introduce the expression

$$\frac{1}{2T} \int_{-T}^T |(U_t g, g)|^2 dt.$$

It has the value

$$\begin{aligned} & \frac{1}{2T} \int_{-T}^T \left| \int_{-\infty}^{+\infty} e^{i\lambda t} d_\lambda \|E(\lambda)g\|^2 \right|^2 dt \\ &= \frac{1}{2T} \int_{-T}^T dt \int_{-\infty}^{+\infty} e^{i\lambda t} d_\lambda \|E(\lambda)g\|^2 \cdot \int_{-\infty}^{+\infty} e^{-i\mu t} d_\mu \|E(\mu)g\|^2 \\ &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \left(\frac{1}{2T} \int_{-T}^T e^{i(\lambda-\mu)t} dt \right) d_\lambda \|E(\lambda)g\|^2 \cdot d_\mu \|E(\mu)g\|^2 \\ &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{e^{i(\lambda-\mu)T} - e^{-i(\lambda-\mu)T}}{2i(\lambda-\mu)T} \cdot d_\lambda \|E(\lambda)g\|^2 \cdot d_\mu \|E(\mu)g\|^2 \\ &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{\sin(\lambda-\mu)T}{(\lambda-\mu)T} \cdot d_\lambda \|E(\lambda)g\|^2 \cdot d_\mu \|E(\mu)g\|^2. \end{aligned}$$

Since $E(\lambda)$ has no point-spectrum except possibly for $\lambda = 0$, for which value precisely we have $(E(+0) - E(-0))g = E_0 g = 0$, $\|E(\lambda)g\|^2$ is a

continuous function of λ , which, as λ goes from $-\infty$ to $+\infty$, goes in a non-decreasing manner from 0 to $\|g\|^2$. If $\lambda = \phi(x)$ is the inverse of $x = \|E(\lambda)g\|^2$, $\phi(x)$ increases monotonically from $-\infty$ to $+\infty$ as x goes from 0 to $\|g\|^2$, it is constant on no interval, although it may have finite jumps. Hence the above expression may be written as

$$\int_0^{\|g\|^2} \int_0^{\|g\|^2} \frac{\sin\{\phi(x) - \phi(y)\} T}{\{\phi(x) - \phi(y)\} T} dx dy.$$

Since the region of integration is finite and the integrand in absolute value less than 1, and since (except for the zero set $x = y$) it approaches zero, the expression approaches zero as $T \rightarrow \pm\infty$. (This would not be true if $\phi(x)$ were constant on certain intervals, i.e., if $\|E(\lambda)g\|^2$ had discontinuities, and thus, $E(\lambda)$ had a point spectrum.)

Thus,

$$\lim_{T \rightarrow +\infty} \frac{1}{2T} \int_{-T}^{+T} |(U_t g, g)|^2 dt = 0. \quad (2)$$

Consider the series

$$\pi(t) = \sum_{n=1}^{\infty} \frac{1}{2^n \|g_n\|^2} |(U_t g_n, g_n)|^2.$$

It is uniformly convergent for $|t| < +\infty$, since it is dominated by the series $\sum_{n=1}^{\infty} \frac{1}{2^n}$, on account of the inequality $|(U_t g_n, g_n)|^2 \leq \|U_t g\| \cdot \|g\| = \|g\|^2$. Hence, in virtue of (2),

$$\lim_{t \rightarrow +\infty} \frac{1}{2T} \int_{-T}^T \pi(t) dt = 0. \quad (3)$$

Now if we can show that a zero t -set I can be found such that

$$\lim_{\substack{t \rightarrow \pm\infty \\ t \text{ not on } I}} \pi(t) = 0, \quad (4)$$

then, on account of the inequality

$$0 \leq |(U_t g_n, g_n)|^2 \leq 2^n \|g\|^2 \pi(t),$$

we shall have the consequence that

$$\lim_{\substack{t \rightarrow \pm\infty \\ t \text{ not on } I}} (U_t g_n, g_n) = 0,$$

that is, the theorem will be proved.

Let $I_m = \left[\pi(t) \geq \frac{1}{m} \right]$. Evidently $I_1 \subset I_2 \subset I_3 \subset \dots$, and, in virtue of (3),

$$\lim_{T \rightarrow +\infty} \frac{m(I_m \cdot [|t| \leq T])}{2T} = 0.$$

We now choose $T_1 < T_2 < \dots$ such that the left-hand member of the above is, for all $m = 2, 3, \dots$, and $T \geq T_{m-1}$, less than $\frac{1}{m}$. Then we set

$$I = I_1 \cdot [|t| \leq T_1] + I_2 \cdot [T_1 < |t| \leq T_2] + \dots \\ \dots + I_m \cdot [T_{m-1} < |t| \leq T_m] + \dots$$

For $T_{m-1} < T \leq T_m$ we have

$$I \cdot [|t| \leq T] \subset I_{m-1} \cdot [|t| \leq T_{m-1}] + I_m \cdot [|t| \leq T],$$

hence,

$$\begin{aligned} \frac{m(I \cdot [|t| \leq T])}{2T} &\leq \frac{m(I_{m-1} \cdot [|t| \leq T_{m-1}]) + m(I_m \cdot [|t| \leq T])}{2T} \\ &\leq \frac{m(I_{m-1} \cdot [|t| \leq T_{m-1}])}{2T_{m-1}} + \frac{m(I_m \cdot [|t| \leq T])}{2T} \\ &\leq \frac{1}{m-1} + \frac{1}{m}. \end{aligned}$$

As $m \rightarrow +\infty$, that is, as $T \rightarrow +\infty$, this approaches zero; hence I is a zero t -set. Since, further, when $|t| > T_{m-1}$, a t on I_m is on I , it follows that outside I , $\pi(t) < \frac{1}{m}$. Hence as $m \rightarrow +\infty$, or $t \rightarrow \pm\infty$, $\pi(t) \rightarrow 0$, so that (4) is established.

Corollary.—The set I of Theorem I can be taken as a set of intervals such that there are only a finite number of intervals of the set in any arbitrarily chosen finite interval.

Proof.—Since the expressions $(U_i g, g)$ are continuous function of t , $\pi(t)$ is likewise, so that every I_n is a closed set, and, likewise, I . Hence the Lebesgue and that Jordan measure of I are equal. In each of the regions $|t| \leq 1, |t| \leq 2, \dots, m-1 < |t| \leq m, \dots$ we can replace I by a finite set of overlapping intervals without increasing its measure by more than, for example, $\frac{1}{2}, \frac{1}{4}, \dots, \frac{1}{2^m}, \dots$,—thus, in all, by a quantity < 1 , so that it remains a zero t -set.

THEOREM II. When the hypothesis C is not realized, the conclusion of Theorem I is invalid.

Proof.—To begin with, suppose the system to be integrable. Then there will exist an invariant $\Lambda \subset \Omega : \Lambda_t = \Lambda$, with $\mu\Lambda > 0$ and finite, and

$\mu(\Omega - \Lambda) > 0$. If $\mu\Omega$ is finite, take $M = \Lambda$, $N = \Omega - \Lambda$ in (1), when $\mu(M_t \cdot N) = \mu(\Lambda_t \cdot (\Omega - \Lambda)) = \mu(\Lambda \cdot (\Omega - \Lambda)) = 0$, whereas $\frac{\mu\Lambda \cdot \mu(\Omega - \Lambda)}{\mu\Omega} \neq 0$. If $\mu\Omega = \infty$, take $M = N = \Lambda$, in which case $\mu(M_t \cdot N) = \mu(\Lambda_t \cdot \Lambda) = \mu\Lambda > 0$, whereas $\frac{\mu\Lambda \cdot \mu\Lambda}{\mu\Omega} = 0$. Thus, in either case, (1) is violated.

Suppose, on the other hand, that the system is non-integrable, but that it has a characteristic function φ of characteristic number $\lambda > 0$: $U_t\varphi = e^{i\lambda t}\varphi$. In this case, as we have seen, we may take $\varphi = e^{i\theta}$, and $\mu\Omega$ will be finite. In (1') take $f = g = \varphi$; then $(U_tf, g) = (e^{i\lambda t}\varphi, \varphi) = e^{i\lambda t}(\varphi, \varphi) = e^{i\lambda t}\mu\Omega$, so that (1'), which in the non-integrable case is a consequence of (1), is impossible.

Theorem II is thus proved.

Now let us consider to what extent the presence of the exceptional t -set I appears necessary in Theorem I. We have:

THEOREM III. There exist spectra $E(\lambda)$ (not necessarily belonging to a dynamical system¹⁰) such that, under the hypothesis C , the conclusion (1') of Theorem I fails to hold when I is empty.

Proof.—Take f such that $E_0f = 0$; then it is sufficient to show that $\lim_{t \rightarrow \pm \infty} (U_tf, f)$ does not exist, or else that it is not zero. We have¹¹

$$(U_tf, f) = \int_{-\infty}^{+\infty} e^{i\lambda t} d\|E(\lambda)f\|^2,$$

and since $x = \|E(\lambda)f\|^2$ can evidently be taken to be an arbitrarily given continuous non-decreasing function, increasing from zero to a finite a ($= \|f\|^2$) as t goes from $-\infty$ to $+\infty$, its inverse $\lambda = \varphi(x)$ can obviously be taken to be equal to an arbitrary monotonically increasing function, possibly with finite jumps, in the interval of definition $0 \leq x \leq a$, and going from $-\infty$ to $+\infty$. Hence

$$(U_tf, f) = \int_0^a e^{i\varphi(x)t} dx.$$

Now let $a = 1$, and $\varphi(x)$ have always such a value that its hexadic development contains only the digits 0 and 5. For instance, let $\varphi(0) = -\infty$, $\varphi(1) = +\infty$, and, for $x = \sum_{n=1}^{\infty} \frac{a_n}{2^n}$ ($a_n = 0$ or 1, the latter infinitely often), let $\varphi(x) = \sum_{n=1}^{\infty} \frac{5a_n}{6^n}$. For $t = 2\pi \cdot 6^n$, the fractional part of $\frac{1}{2\pi} \varphi(x) \cdot t$ will thus start with the hexadic digit 0 or 5, i.e., it is ≥ 0 , $\leq \frac{1}{6}$, or $\geq \frac{5}{6}$.

≤ 1 —i.e., $\frac{1}{2\pi} \varphi(x)t$ has a value (mod 1) $\geq -\frac{1}{6}$, $\leq \frac{1}{6}$. Hence

$$\begin{aligned} \Re \int_0^a e^{i\varphi(x)t} dx &= \int_0^1 \cos \left(2\pi \cdot \frac{1}{2\pi} \varphi(x)t \right) dx \\ &\geq \int_0^1 \cos \left(2\pi \cdot \frac{1}{6} \right) dx \\ &= \cos \frac{\pi}{3} = \frac{1}{2} > 0. \end{aligned}$$

Thus $\lim_{t \rightarrow \pm \infty} (U_t f, f) = 0$ is excluded, and our theorem is proved.

III. Theorems I, II, III, delimit a situation which could scarcely be rendered more precise: the hypotheses and diverse restrictions seem all inherent in the nature of the case. It would, of course, be extremely interesting if a proof of Theorem I could be found without the use of the spectrum $E(\lambda)$, etc., for example, along the lines of the recent paper of Birkhoff's on the ergodic hypothesis.¹²

Theorem I expresses the fact that in a system for which C holds, consecutive states at two sufficiently different epochs are almost certainly statistically independent. For $\frac{\mu(M \cdot N)}{\mu\Omega}$ is the probability that, in the

time t , the system go from M to N —and this is $= \frac{\mu M \cdot \mu N}{(\mu\Omega)^2}$, i.e., the product of the probabilities of the system's being in M and of its being in N . (We say "almost certainly," to correspond with the fact that the exceptional zero t -set I is excluded.) Such a system can have, *a la longue*, no physical properties. Indeed, the only properties which any system can have *a la longue* are those which violate C , namely:

1. Invariant sub-sets: barriers that are never passed.
2. Angle variables: clocks that never change.¹³

Theorem I may also be expressed by saying that the states of motion corresponding to any set M of Ω become more and more spread out into an amorphous everywhere dense chaos. Periodic orbits, and such like, appear only as very special possibilities of negligible probability.

We have already called attention to the essential difference between this situation, which may exist in very simple systems, and that envisaged in the kinetic theory of gases, in which the confused character of the motion is an intuitively evident consequence of the large number of degrees of freedom of the system.

¹ "Hamiltonian Systems and Transformations in Hilbert Space," these PROCEEDINGS, 315–318 (May, 1931); this reference will here be abbreviated to (H). Another reference of importance for us is the paper by v. Neumann on the proof of the quasi-ergodic hy-

pothesis (these PROCEEDINGS, pp. 70-82 (Jan., 1932)), which will be referred to under the abbreviation (Q).

We shall, in the present paper, assume the notation, etc., of (H) and (Q) to be known.

² Properties, of course, which are true "almost everywhere" on the manifold of states of motion Ω , corresponding with the nature of the functional tools. The allied conception of properties true "almost everywhere" on the time axis with respect to the time-average τ -measure here defined is an essential notion in the present paper; it is contained implicitly in J. v. Neumann's recent proof of the quasi-ergodic hypothesis (these PROCEEDINGS, pp. 70-82 (Jan., 1932), and 263-266 (March, 1932)).

³ These cases are made the subject of as yet unpublished investigations by B. O. Koopman. In the case of a pure point spectrum, the theory of almost periodic functions is available: $(U_t f, g)$ is an almost periodic function of t , a fact having an obvious physical interpretation when $f = \chi_M(P)$ and $g = \chi_N(P)$, characteristic functions of the sets M and N ($\subset \Omega$; $\mu M, \mu N$ finite). In the second case here noted, in virtue of the well-known Fourier integral theorem, $(U_t f, g) \rightarrow 0$ as $t \rightarrow \infty$, showing (on making the above choice of f, g) that any such region M will flow, in course of time, "almost entirely" out of a fixed region N of finite measure. This possibility obviously implies that $\mu\Omega$ is infinite. There is a similarly obvious interpretation when $\lambda=0$ is a simple unique characteristic number.

⁴ Cf. (H), p. 318.

⁵ It is unnecessary for us to enter upon the discussion of τ -measurability or the properties of τ -measure.

⁶ Cf. (Q), p. 78.

⁷ The ergodic theorem (Q) states that $\frac{1}{T} \int_0^T U_t f dt$ converges "strongly" to $E_0 f$, thus this theorem tells us *more* than the present one from the point of view of convergence, but *less*, from the point of view of the function, since f is replaced by its time average.

⁸ If we replace in $(U_t f, f) \rightarrow (E_0 f, f)$ f by $\frac{f+g}{2}$ and by $\frac{f-g}{2}$, and subtract, the real part of $(U_t f, g) \rightarrow (E_0 f, g)$ results. If we replace g by ig , the imaginary part results. Thus the general formula holds.

⁹ Cf. (Q), p. 73.

¹⁰ The following question is of great interest: When will a canonical resolution of the identity $E(\lambda)$ be the spectrum of a dynamical system? An (unpublished) formula obtained by Koopman is that, when f, g , and the ordinary product fg are in $\mathfrak{S}$, then, for a dynamical system,

$$E(x)fg = \int_{-\infty}^{+\infty} E(x-s)f \cdot d_s E(s)g,$$

which is a consequence of the fundamental equation obtained in (H):

$$U_t F(f, g, \dots) = F(U_t f, U_t g, \dots),$$

where F is any single-valued function. Recently J. v. Neumann has shown that the first-mentioned formula is sufficient and necessary for the U_t being generated by a suitably chosen group of point-transformations $P \rightarrow P_t$ (unpublished).

¹¹ If $\| \mathfrak{L}(\lambda)f \|^2$ is differentiable with respect to λ (the case for a type of continuous spectrum defined by Hellinger and Hahn) we have $\int_{-\infty}^{+\infty} \frac{d}{d\lambda} \| E(\lambda)f \|^2 \cdot e^{it\lambda} d\lambda$. By a well-known theorem on Fourier integrals this has $\lim$ equal to zero. Cf. reference 3 above.

$t \rightarrow \infty$

¹² These PROCEEDINGS, pp. 650–660 (Dec., 1931).

¹³ Though it is very probable that the case C or C' , in which our theorems hold is the general one for dynamical systems, it is not easy to construct effective examples. (See footnote 10.) An example, recently constructed by v. Neumann, will be published soon; it refers to a two-dimensional flow of the following type: the flow takes place in a rectangle, oriented parallel to the X and Y axes, the upper side of which has been replaced by suitable chosen curve $Y = F(X)$. The flow itself is parallel to the positive Y -axis, and each point X , $F(X)$ has to be identified with the corresponding point $X + \alpha$, 0. (The number $X + \alpha$ is to be taken mod. a , where a is the breadth of the parallelogram in the direction of the X -axis; α is a number incommensurable with a .) If $F(X)$ and α are suitably chosen this flow can be shown to fulfill C (and C').

ÜBER EINEN SATZ VON HERRN M. H. STONE

1. Herr M. H. Stone hat vor einiger Zeit den folgenden Satz formuliert und bewiesen²:

(S) Sei U_t eine Schar unitärer Operatoren des Hilbertschen Raumes, mit der Gruppeneigenschaft

$$U_t U_s = U_{t+s}.$$

(t, s reelle Zahlen). Dann existiert eine Zerlegung der Einheit $E(\lambda)$ ³, so daß symbolisch

$$U_t = \int_{-\infty}^{+\infty} e^{it\lambda} dE(\lambda)$$

gilt⁴.

Hierbei ist Voraussetzung, daß U_t stetig von t abhängt⁵. Die Anwendungen, die von diesem Satze, besonders in der von Herrn B. O. Koopman entdeckten Operatorenbehandlungsweise der klassischen Mechanik gemacht werden können⁶ (so konnte der Verfasser mit ihrer Hilfe die klassisch-mechanische Quasi-Ergodenhypothese beweisen⁷), lassen es als erwünschenswert erscheinen, den Stoneschen Satz auch ohne die Stetigkeitsannahme zu beweisen⁸. Auch rein mathematisch ist eine solche Verschärfung anzustreben.

In einer früheren Abhandlung⁹ führte der Verfasser den Begriff der meßbaren Abhängigkeit eines Operators A_t vom Zahlenparameter t ein:

² Proc. Nat. Ac. 16, February 1930, S. 173-174.

³ Vgl. z. B. die Abhandlung des Verfassers „Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren“, Math. Ann. 102,1 (1929). Sie soll im folgenden als E zitiert werden.

⁴ D. h. für zwei beliebige f, g des Hilbertschen Raumes

$$(U_t f, g) = \int_{-\infty}^{+\infty} e^{it\lambda} d(E(\lambda) f, g).$$

Die Bezeichnung nach a. O. Anmerkung ², sowie E.

⁵ Vgl. z. B. die Abhandlung des Verfassers „Zur Algebra der Funktionaloperatoren usw.“, Math. Ann. 102,3 (1929), wo der Stetigkeitsbegriff variabler Operatoren präzisiert wird, insbesondere S. 381-388. Diese Arbeit soll im folgenden als A zitiert werden.

⁶ Proc. Nat. Ac. 17, May 1931.

⁷ Proc. Nat. Ac. 18, January 1932, sowie weitere Abhandlungen in den Proceedings, und den zweitfolgenden Artikel dieses Annals-Heftes.

⁸ Vgl. den in Anmerkung ⁷ genannten Artikel des Verfassers in diesem Annals-Bande.

⁹ Math. Ann. 104,4 (1931), S. 572.

(M) Der Operator A_t hängt meßbar vom Zahlenparameter t ab, wenn für alle f, g des Hilbertschen Raumes die numerische Funktion $(A_t f, g)$ im gewöhnlichen (Lebesgueschen) Sinne meßbar ist. Übrigens kann man bei dieser Definition f, g auf die Elemente eines vollständigen normierten Orthogonalsystems $\varphi_1, \varphi_2, \dots$ beschränken¹⁰.

Im nächsten Paragraphen zeigen wir, daß die Prämissen von (S) (d. i.: U_t unitär, $U_t U_s = U_{t+s}$) zusammen mit (M) (d. i.: U_t meßbar von t abhängig) die Stetigkeit von U_t in bezug auf t nach sich zieht. Im letzten Paragraphen geben wir einen neuen Beweis des Stoneschen Satzes an.

2. Erfülle U_t die Prämisse von (S) sowie (M). Da $(U_t f, g)$ meßbar von t abhängt, (f, g beliebig, aber vorläufig fest im Hilbertschen Raume), gibt es nach einem Satze von Lusin¹¹ zu jedem $\varepsilon > 0$ eine t -Menge $\mathfrak{T} = \mathfrak{T}_{f,g,\varepsilon}$ von Lebesgueschem Maße $< \varepsilon$, so daß $(U_t f, g)$, als Funktion des auf die Komplementärmenge von $\mathfrak{T}$ beschränkten t betrachtet, stetig ist. Sei $f_1, f_2, \dots$ eine im Hilbertschen Raume überall dichte Folge¹², durchlaufe f, g alle Paare f_m, f_n . Wir setzen dazu $\varepsilon = \frac{\delta}{2^{m+n}}$ ($\delta > 0$). Die Menge

$\mathfrak{T}' = \mathfrak{T}'_\delta = \bigcap_1^{\infty} \bigcap_{m,n} \mathfrak{T}_{f_m, f_n, \frac{\delta}{2^{m+n}}}$ hat ein Maß $< \sum_1^{\infty} \frac{\delta}{2^{m+n}} = \delta$, und für das

auf ihre Komplementärmenge beschränkte t sind alle $(U_t f_m, f_n)$ stetig, also auch alle $(U_t f, g)$ (denn $(U_t f, g)$ ist, wenn t als Parameter angesehen wird, gleichmäßig stetig in f, g für alle t ¹³). Wenn also t und $t_1, t_2, \dots$ nicht in $\mathfrak{T}'$ liegen, und $t_n \rightarrow t$, so gilt für alle f, g $(U_{t_n} f, g) \rightarrow (U_t f, g)$, d. h. die $U_{t_n} f$ konvergieren schwach¹⁴ gegen $U_t f$. Da aber gleichzeitig $\|U_{t_n} f\|$ gegen $\|U_t f\|$ konvergiert (wegen der Unitarität sind ja alle $\|f\|$), konvergiert $U_{t_n} f$ sogar stark gegen $U_t f$ ¹⁵. Nach einem bekannten Satze von Lebesgue hat in allen Punkten t außerhalb von $\mathfrak{T}'$, bis auf eine Menge vom Maße 0, dieses $\mathfrak{T}'$ die Dichte 0¹⁶. Für ein solches t können

¹⁰ D. h.: die Matricelemente der A_t (im Koordinatensystem der $\varphi_1, \varphi_2, \dots$) sollen meßbare Funktionen von t sein. Die Gleichwertigkeit der Definitionen beweist man leicht, vgl. a. a. O. Anmerkung ⁹.

¹¹ C. R., 154, 1912, S. 1688, oder Sierpinski, Fund. Math. 3, 1922, S. 320.

¹² Der Hilbertsche Raum ist separabel! Vgl. z. B. E, S. 65 und 109–111.

¹³ Da es in f, g linear ist, genügt es, die Fälle $f = 0$ oder $g = 0$ zu betrachten. Dann ist $(U_t f, g) = 0$, und wegen der Schwarzschen Ungleichheit sowie der Unitarität von U_t

$$|(U_t f, g)| \leq \|U_t f\| \cdot \|g\| = \|f\| \cdot \|g\|.$$

Vgl. E, S. 64, Satz 1.

¹⁴ Vgl. z. B. A, S. 378–381.

¹⁵ Wenn f_n schwach gegen f konvergiert, und $\|f_n\|$ (numerisch) gegen $\|f\|$, so gilt sogar starke Konvergenz:

$$\|f_n - f\|^2 = \|f_n\|^2 + \|f\|^2 - 2\Re(f, f_n) \rightarrow 2\|f\|^2 - 2\Re(f, f) = 0.$$

¹⁶ Vgl. z. B. Carathéodory, „Reelle Funktionen“, Berlin 1918, S. 492–497, Satz 3. Die dort auftretende Funktion $f(P)$ sei $f(t)$, u. zw. $= 1$ in $\mathfrak{T}'$, sonst $= 0$.

wir also zu jedem f , $\theta > 0$, $\eta > 0$ folgendes finden: Erstens ein $\delta_1 > 0$, so daß für $\delta' \leq \delta_1$, > 0 der im Intervalle $t - \delta' < t' < t + \delta'$ liegende Teil von $\mathfrak{T}'$ ein Maß $< \theta \cdot \delta'$ hat (wegen der Dichte 0!). Zweitens ein $\delta_2 > 0$, so daß für jedes nicht zu $\mathfrak{T}'$ gehörige t' des Intervalles $t - \delta_2 < t' < t + \delta_2$ $\|U_{t'} f - U_t f\| < \eta$ gilt (wegen der Stetigkeit). Sei $\delta_0 = \min(\delta_1, \delta_2) > 0$, also $\delta_0 = \delta_0(t, f, \theta, \eta)$. Für $\delta' \leq \delta_0$, > 0 gilt also im Intervalle $t - \delta' < t' < t + \delta'$ überall, bis auf eine Menge vom Maße $< \theta \cdot \delta'$, $\|U_{t'} f - U_t f\| < \eta$. In dieser letzteren Aussage kommen $\mathfrak{T}' = \mathfrak{T}'_{\delta'}$, δ nicht mehr vor, sondern nur U_t und t, f, θ, η sowie δ_0 ; wir wollen sie, die analoge Denjoysche Begriffsbildung verallgemeinernd, approximative starke Stetigkeit von U_t in t nennen (vgl. a. a. O. Anm. ¹¹). Sie gilt für alle t außerhalb von $\mathfrak{T}'$, bis auf eine Menge vom Maße 0, d. h. für alle t bis auf eine Menge vom Maße $< \delta$. Da aber $\delta > 0$ beliebig war, für alle t , bis auf eine Menge vom Maße 0.

Wegen $U_t = U_{t-t_0} U_{t_0}$, $U_{t'} = U_{t-t_0} U_{t_0+(t'-t)}$ und der Unitarität von U_{t-t_0} ist U_t überall approximativ stark stetig, wenn es dies für t_0 ist — also ist ersteres der Fall.

Sei nun ein f und $\epsilon > 0$ gegeben, wir wählen ein $\delta > 0$ so, daß für alle t' in $-\delta < t' < \delta$, bis auf eine Menge vom Maße $< \frac{1}{2} \delta$, $\|U_{t'} f - f\| < \frac{\epsilon}{2}$ ist ($t = 0$, $\theta = \frac{1}{2}$, $\eta = \frac{\epsilon}{2}$). Ferner seien t_1, t_2 beliebig, aber mit $|t_1 - t_2| < \delta$. Für $\left| \tau - \frac{t_1 + t_2}{2} \right| < \frac{\delta}{2}$ ist bestimmt $|t_1 - \tau| < \delta$, $|\tau - t_2| < \delta$, also gilt auch hier bis auf Mengen von Maßen $< \frac{\delta}{2}$, $\|U_{t_1-\tau} f - f\| < \frac{\epsilon}{2}$, $\|U_{\tau-t_2} f - f\| < \frac{\epsilon}{2}$. Bis auf eine Menge vom Maße $< \delta$ gilt daher beides zugleich, und da das Intervall $\left| \tau - \frac{t_1 + t_2}{2} \right| < \frac{\delta}{2}$ das Maß δ hat, gibt es ein τ mit allen diesen Eigenschaften. Hieraus folgt aber:

$$\begin{aligned} \|U_{t_1} f - U_{t_2} f\| &\leq \|U_{t_1} f - U_{\tau} f\| + \|U_{\tau} f - U_{t_2} f\| \\ &= \|U_{\tau}(U_{t_1-\tau} f - f)\| + \|U_{t_2}(U_{\tau-t_2} f - f)\| \\ &= \|U_{t_1-\tau} f - f\| + \|U_{\tau-t_2} f - f\| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon. \end{aligned}$$

Damit ist die starke Stetigkeit von U_t für alle t (vgl. a. a. O. Anm. ⁵) bewiesen.

3. Wir gehen nun daran, den Stoneschen Satz zu beweisen.

Sei $a(t)$ eine numerische Funktion mit endlichem $\int_{-\infty}^{+\infty} |a(t)| dt$, dann

können wir einen Operator A symbolisch durch $\int_{-\infty}^{+\infty} a(t) U_t dt$ definieren, d. h. durch

$$(Af, g) = \int_{-\infty}^{+\infty} a(t) (U_t f, g) dt$$

(vgl. a. a. O. Anm. ⁹, S. 573–574). Man beweist leicht, wenn auch $\int_{-\infty}^{+\infty} |b(t)| dt$ endlich und $B = \int_{-\infty}^{+\infty} b(t) U_t dt$ ist, die folgenden Formeln:

$$A^* = \int_{-\infty}^{+\infty} \overline{a(-t)} U_t dt, \quad cA = \int_{-\infty}^{+\infty} ca(t) U_t dt,$$

$$A + B = \int_{-\infty}^{+\infty} (a(t) + b(t)) U_t dt,$$

$$U_s A = A U_s = \int_{-\infty}^{+\infty} a(t-s) U_t dt,$$

$$AB = \int_{-\infty}^{+\infty} c(t) U_t dt, \quad \left(c(t) = \int_{-\infty}^{+\infty} a(s) b(t-s) ds \right)^{17}$$

(vgl. die analogen Beweise am oben a. O.). Die letzteren Formeln zeigen, daß unsere $A, B, \dots$ miteinander und mit den U_t vertauschbar sind.

Wir setzen nun $a(t) = \begin{cases} e^{-t} & \text{für } t \geq 0 \\ 0 & \text{für } t < 0 \end{cases}$, d. h.

$$A = \int_0^{+\infty} e^{-t} U_t dt, \quad A^* = \int_{-\infty}^0 e^t U_t dt.$$

Daraus folgt sofort:

$$\left. \begin{aligned} A + A^* &= \int_{-\infty}^{+\infty} e^{-|t|} U_t dt, \\ AA^* &= A^*A = \int_{-\infty}^{+\infty} \frac{1}{2} e^{-|t|} U_t dt, \end{aligned} \right\} AA^* = A^*A = \frac{1}{2}(A + A^*).$$

Für $V = 1 - 2A$ ergibt das: $VV^* = V^*V = 1$, d. h. V ist unitär. Es ist, ebenso wie A , mit den U_t vertauschbar.

Betrachten wir den Ausdruck $\frac{U_t - 1}{t} (V - 1)$. Er ist gleich

$$-\frac{2}{t} (U_t - 1)A = \frac{2}{t} (A - U_t A),$$

also gleich

$$\begin{aligned} & \frac{2}{t} \left(\int_0^\infty e^{-s} U_s ds - \int_t^\infty e^{-s+t} U_s ds \right) \\ &= 2e^t \cdot \frac{1}{t} \int_0^t e^{-s} U_s ds - 2 \frac{e^t - 1}{t} \cdot \int_0^\infty e^{-s} U_s ds. \end{aligned}$$

¹⁷ A^* ist die Adjungierte von A . Man beachte, daß aus $U_t U_s = U_{t+s}$ folgt. $U_0 = 1$, $U_t^{-1} = U_{-t}$, als wegen der Unitarität $U_t^* = U_t^{-1} = U_{-t}$.

Im ersten Glied konvergiert der erste Faktor für $t \rightarrow 0$ gegen 2, der zweite (der Operator!) wegen der Stetigkeit von U_s gegen $U_0 = 1$ (stark!), im zweiten Glied konvergiert der erste Faktor gegen 2, der zweite (der Operator!) ist gleich A . Also ist der (starke Operatoren-)Limes gleich $2 \cdot 1 - 2 \cdot A = V + 1$. Etwas anders geschrieben:

$$\frac{U_t - 1}{it} (V - 1) \xrightarrow{\text{stark}} -i(V + 1) \quad (t \rightarrow 0).$$

Aus dieser Formel folgt Verschiedenes. Erstens hat $Vf = f$ $(V - 1)f = 0$ zur Folge, also

$$\frac{U_t - 1}{it} (V - 1)f = 0, \quad i(V + 1)f = 0, \quad \text{d. h. } Vf = -f.$$

Also $f = 0$. Somit ist V die Cayleysche Transformierte eines hypermaximalen Operators R .¹⁸ Rg hat Sinn, wenn $g = (V - 1)f$ ist, und dann ist es gleich $-i(V + 1)f$,¹⁸ also gilt für alle g , für die Rg Sinn hat,

$$\frac{U_t - 1}{it} g \xrightarrow{\text{stark}} Rg \quad (t \rightarrow 0).$$

Wenn $\frac{U_t - 1}{it} g$ auch nur einen schwachen Limes für $t \rightarrow 0$ hat, können wir diesen gleich $R'g$ setzen, dadurch einen Operator R' definierend. Aus

$$\left(\frac{U_t - 1}{it} g, h \right) = \left(g, \frac{U_t^* - 1}{-it} h \right) = \left(g, \frac{U_{-t} - 1}{i(-t)} h \right)$$

folgt $(R'g, h) = (g, R'h)$, so daß R' Hermitesch ist. Da es nach Obigem das hypermaximale R fortsetzt, ist $R' = R$. D. h.: wenn Rg keinen Sinn hat, konvergiert $\frac{U_t - 1}{it} g$ nicht einmal schwach.

Sei schließlich $E(\lambda)$ die zu R gehörige Zerlegung der Einheit.¹⁹ Wir definieren eine Operatorenschar U'_t symbolisch durch $\int_{-\infty}^{+\infty} e^{it\lambda} dE(\lambda)$, d. h. durch

$$(U'_t f, g) = \int_{-\infty}^{+\infty} e^{it\lambda} d(E(\lambda) f, g),$$

ganz analog wie in einer früheren Abhandlung des Verfassers.²⁰ Durch dieselben Rechnungen wie dort²¹ erkennt man:

¹⁸ E, S. 80–81 und 91.

¹⁹ E, S. 91, Def. 17 und S. 92, Satz 36.

²⁰ Annals of Math., 32/2 (1931).

²¹ Vgl. insbesondere S. 205, a)–b), und deren Beweise.

$$U'_0 = \int_{-\infty}^{+\infty} dE(\lambda) = 1,$$

$$U'_t{}^* = \int_{-\infty}^{+\infty} e^{it\lambda} dE(\lambda) = \int_{-\infty}^{+\infty} e^{-it\lambda} dE(\lambda) = U'_{-t},$$

$$U'_t U'_s = \int_{-\infty}^{+\infty} e^{it\lambda} \cdot e^{is\lambda} dE(\lambda) = \int_{-\infty}^{+\infty} e^{i(t+s)\lambda} dE(\lambda) = U'_{t+s}.$$

Also ist $U'_t U'_t{}^* = U'_t U'_{-t} = U'_0 = 1$, ebenso $U'_t{}^* U'_t = 1$, d. h. U'_t unitär, mit derselben Gruppeneigenschaft wie die U_t . Ferner ist jedes U_t mit V vertauschbar, also mit allen $E(\lambda)$,²² also auch mit allen U'_s .

Wenn für $\lambda \geq C$ $E(\lambda)f = f$ ist, und für $\lambda \leq -C$ $E(\lambda)f = 0$ (was z. B. für jedes $f = (E(C) - E(-C))g$ der Fall ist²³), so können wir bei $U'_t f$ das $\int_{-\infty}^{+\infty}$ durch ein $\int_{-C}^{+C}$ ersetzen. Dann erkennt man mühelos die Legitimität der folgenden Rechnung (das $\xrightarrow{\text{stark}}$ gilt für $t \rightarrow 0$):

$$\begin{aligned} \frac{U'_t - 1}{it} f &= \frac{\int_{-C}^{+C} e^{it\lambda} dE(\lambda) f - \int_{-C}^{+C} dE(\lambda) f}{it} \\ &= \int_{-C}^{+C} \frac{e^{it\lambda} - 1}{it} dE(\lambda) f \xrightarrow{\text{stark}} \int_{-C}^{+C} \lambda dE(\lambda) f \\ &= \int_{-\infty}^{+\infty} \lambda dE(\lambda) f = Rf. \end{aligned}$$

Für diese f hat Rf Sinn, und sie liegen überall dicht, da für $C \rightarrow +\infty$ $(E(C) - E(-C))g \rightarrow g = 0$.²² Somit gilt in einer überall dichten

Menge der f : $\frac{U'_t - 1}{it} f \rightarrow Rf$, $\frac{U_t - 1}{it} f \rightarrow Rf$, also

$$i \left(\frac{U_t - 1}{it} f - \frac{U'_t - 1}{it} f \right) = \frac{U_t - U'_t}{t} f \rightarrow 0.$$

Für ein solches f ist aber:

$$\begin{aligned} \|(U_t - U'_t)f\| &= \left\| \sum_1^n \left(U_{\frac{m+1-r}{n}t} U'_{\frac{r-1}{n}t} - U_{\frac{n-r}{n}t} U'_{\frac{r}{n}t} \right) f \right\| \\ &= \left\| \sum_1^n U_{\frac{n-r}{n}t} U'_{\frac{r-1}{n}t} \left(\frac{U_t}{n} - \frac{U'_t}{n} \right) f \right\| \\ &\leq \sum_1^n \left\| U_{\frac{n-r}{n}t} U'_{\frac{r-1}{n}t} \left(\frac{U_t}{n} - \frac{U'_t}{n} \right) f \right\| \end{aligned}$$

²² Vgl. z. B. E, S. 115, Satz 10* und dessen Beweis, woraus dies leicht folgt.

²³ Vgl. die Formeln in E, S. 91, Def. 17.

$$\begin{aligned}
&= \sum_1^n \left\| \left(U_{\frac{t}{n}} - U'_{\frac{t}{n}} \right) f \right\| \\
&= n \left\| \left(U_{\frac{t}{n}} - U'_{\frac{t}{n}} \right) f \right\| \\
&= t \cdot \left\| \frac{U_{\frac{t}{n}} - U'_{\frac{t}{n}}}{\frac{t}{n}} f \right\|.
\end{aligned}$$

Für $n \rightarrow +\infty$ strebt der zweite Faktor gegen 0, also ist

$$\|(U_t - U'_t)f\| = 0, \quad U_t f = U'_t f.$$

Da $U_t f$, $U'_t f$ stetig von f abhängen, gilt dies nicht nur in der genannten überall dichten f -Menge, sondern für alle f . Also ist

$$U_t = U'_t = \int_{-\infty}^{+\infty} e^{it\lambda} dE(\lambda).$$

Damit ist die Behauptung von (S) bewiesen.

EINIGE SÄTZE ÜBER MESSBARE ABBILDUNGEN

1. Den Gegenstand dieser Arbeit bilden einige Sätze über meßbare Funktionen, die teilweise wegen ihrer Anwendung im unmittelbar nachfolgenden Artikel des Verfassers abgeleitet wurden. Immerhin sind sie vielleicht auch von selbständigem mathematischen Interesse.

Wir werden im folgenden gewisse Mengen Ω betrachten, die kurz R -räume genannt werden sollen; sie sind so definiert:

DEFINITION 1. Die Menge Ω heißt ein m -Raum, wenn sie metrisch, separabel und vollständig ist² (ihre Elemente sollen $x, y, \dots$ heißen, die Entfernung $\overline{x, y}$).

Beispiele von m -Räumen sind zahlreich: Der k -dimensionale Euklidische Raum, in dem die Distanz zweier Punkte $\{\xi_1, \dots, \xi_k\}, \{\eta_1, \dots, \eta_k\}$ (die ξ, η sind reelle Zahlen!) durch $\sqrt{(\xi_1 - \eta_1)^2 + \dots + (\xi_k - \eta_k)^2}$ definiert wird, oder auch durch $\text{Max}[|\xi_1 - \eta_1|, \dots, |\xi_k - \eta_k|]$; der Hilbertsche Raum (Punkte: $\{\xi_1, \xi_2, \dots\}, \xi_1^2 + \xi_2^2 + \dots$ endlich), in dem die Distanz zweier Punkte $\{\xi_1, \xi_2, \dots\}, \{\eta_1, \eta_2, \dots\}$ durch $\sqrt{(\xi_1 - \eta_1)^2 + (\xi_2 - \eta_2)^2 + \dots}$ definiert wird; der ∞ -dimensionale Raum³ (Punkte: alle $\{\xi_1, \xi_2, \dots\}$), in dem die Distanz zweier Punkte $\{\xi_1, \xi_2, \dots\}, \{\eta_1, \eta_2, \dots\}$ durch

$\text{Min}_{n=1,2,\dots} \left\{ \text{Max} \left[|\xi_1 - \eta_1|, \dots, |\xi_n - \eta_n|, \frac{1}{n} \right] \right\}$ definiert wird⁴; jede abgeschlossene Teilmenge eines m -Raumes⁵; usw.

Der zweite Begriff, den wir einführen, ist derjenige eines dem Lebesgueschen äußeren Maß analogen Maßes, definiert in einem m -Raume. Wir nennen es ein l -Maß⁶:

² Vgl. z. B. Hausdorff „Mengenlehre“ (Berlin und Leipzig, 1927), S. 94, 125 und 103.

³ Im ∞ -dimensionalen bestehen, wie man sieht, mehrere, voneinander wesentlich verschiedene, Analoga des Euklidischen Raumes!

⁴ Diese Entfernungsdefinition hat die Wirkung, daß die Folge $\{\xi_1^{(p)}, \xi_2^{(p)}, \dots\}$ ($p = 1, 2, \dots$) dann und nur dann für $p \rightarrow \infty$ gegen $\{\xi_1, \xi_2, \dots\}$ konvergiert, wenn für jedes n $\xi_n^{(p)} \rightarrow \xi_n$ ($p \rightarrow \infty$) gilt.

⁵ Mit derselben Metrik!

⁶ Die nachfolgenden Begriffsbildungen lehnen an jene von Carathéodorys Theorie der regulären Maßfunktionen an, jedoch mit einigen Zusatzbedingungen (d. s. Verschärfungen in I. und V.). Vgl. Carathéodory „Reelle Funktionen“ (Berlin und Leipzig 1918), S. 237—274.

DEFINITION 2. Im m -Raume Ω ist ein l -Maß definiert, wenn jeder Teilmenge N von Ω eine Zahl $\mu^*(N)$ zugeordnet ist, mit den folgenden Eigenschaften:

- I. $\mu^*(N)$ ist 0, oder positiv-endlich, oder $+\infty$. Für jede Kugel ist es endlich. (Eine Kugel mit dem Mittelpunkt x_0 und dem Radius $\varrho > 0$ ist die Menge der x mit $\overline{x_0, x} \leq \varrho$.)
- II. Aus $M \subset N^7$ folgt $\mu^*(M) \leq \mu^*(N)$.
- III. Seien $N_1, N_2, \dots$ endlich oder abzählbar viele Mengen, M ihre Vereinigungsmenge. Dann ist: $\mu^*(M) \leq \mu^*(N_1) + \mu^*(N_2) + \dots$.
- IV. Ist die Entfernung der Mengen $M, N^8 > 0$, so ist $\mu^*(M+N) = \mu^*(M) + \mu^*(N)$.
- V. $\mu^*(M)$ ist die untere Grenze aller $\mu^*(O)$, wenn O sämtliche offenen Mengen $\supset M$ durchläuft⁹.

Wir können nun, in beinahe wörtlicher Wiederholung von Carathéodorys Begriffsbildungen und Beweisen (a. a. O. Anmerkung ⁶) zunächst definieren:

DEFINITION 3. Die Teilmenge M von Ω heißt meßbar, wenn für jede Teilmenge N von Ω

$$\mu^*(N) = \mu^*(MN) + \mu^*(N - MN)^{10}$$

gilt. Statt $\mu^*(M)$ schreiben wir dann $\mu(M)$ und beweisen dann:

- a) jedes offene M ist meßbar.
- b) Sind $M_1, M_2, \dots$ endlich oder abzählbar viele meßbare Mengen, so ist $M_1 + M_2 + \dots$ und $M_1 \cdot M_2 \cdot \dots$ meßbar. Ferner ist mit M, N auch $M - N$ meßbar.
- c) Alle M mit $\mu^*(M) = 0$ sind meßbar.

⁷ D. h.: M ist Teilmenge von N .

⁸ D. h. die untere Grenze aller $\overline{x, y}$, x in M , y in N .

⁹ II. folgt aus V., wir wollen aber die Carathéodorysche Anordnung nicht ändern. — Übrigens könnten wir in V. an Stelle der offenen Mengen $\supset M$ ebensogut alle Borelschen Mengen $\supset M$ treten lassen. Da jede offene Menge Borelsch ist, ist bloß zu zeigen: wenn P Borelsch ist, so existiert ein offenes $O \supset P$ mit $\mu^*(O) < \mu^*(P) + \varepsilon$ ($\varepsilon > 0$). Bilden wir hierzu die Borelschen Klassen $F^{(\xi)}$ nach Hausdorff, S. 178, aber so, daß $F^{(0)}$ aus den offenen Mengen besteht. Ein Borelsches O obiger Art existiert: etwa $O = P$, wählen wir es aus $F^{(\xi)}$ mit kleinstmöglicher Ordnungszahl ξ . Wäre ξ gerade, $\neq 0$, so wäre $O = O_1 \cdot O_2 \cdot \dots, O_n$ aus $F^{(\xi_n)}$, $\xi_n < \xi$, und nach e) im Text $\lim_{n \rightarrow \infty} \mu^*(O_n) = \mu^*(O)$, also ein $\mu^*(O_n) < \mu^*(P) + \varepsilon$, entgegen der Annahme. Wäre ξ ungerade, also $= \eta + 1$, so ist $O = O_1 + O_2 + \dots, O_n$ in $F^{(\eta)}$. Da bei unserer Definition $F^{(1)} = F^{(0)}$ ist, ist $\eta \neq 0$, also $O_n = O_{n1} \cdot O_{n2} \cdot \dots, O_{nm}$ aus $F^{(\eta_{nm})}$, $\eta_{nm} < \eta$. Also ist $\lim_{n \rightarrow \infty} \mu^*(O_{nm}) = \mu^*(O_n)$, und wenn wir $m = m_n$ mit $\mu^*(O_{nm}) < \mu^*(O_n) + \frac{\delta}{2^n}$ ($\delta > 0$) wählen, für $O' = O_{1m_1} + O_{2m_2} + \dots$ $\mu^*(O') < \mu^*(O) + \delta = \mu^*(P) + \varepsilon$, wenn wir $\delta = \mu^*(P) + \varepsilon - \mu^*(O)$ wählen. Dabei gehört O' zu $F^{(\eta)}$, entgegen der Annahme. Also muß $\xi = 0$ sein, d. h. O offen.

¹⁰ $M+N, MN$ oder $M \cdot N, M - N$ bezeichnen die Vereinigungs-, Durchschnitts-, Differenzmenge. Ebenso für mehrere Addenden bzw. Faktoren.

Ferner die bekannten Additivitäts- und Limeseigenschaften des Lebesgueschen Maßes:

- d) Sind $M_1, M_2, \dots$ wie in b) und paarweise elementfremd, so ist

$$\mu(M_1 + M_2 + \dots) = \mu(M_1) + \mu(M_2) + \dots$$
- e) Sind $M_1, M_2, \dots$ wie in b) und $M_1 \supset M_2 \supset \dots$ oder $M_1 \subset M_2 \subset \dots$ und N ihr Limes (d. i. $M_1 \cdot M_2 \cdot \dots$ bzw. $M_1 + M_2 + \dots$), so ist

$$\mu(N) = \lim_{n \rightarrow \infty} \mu(M_n).$$

Aus a), b) folgt insbesondere, daß alle Borelschen Mengen meßbar sind.

Beispiele von l -Mäßen sind leicht anzugeben: das gewöhnliche Lebesguesche äußere Maß im k -dimensionalen Euklidischen Raume; wenn $f(x_1, \dots, x_k)$ eine in diesem Raume definierte, im gewöhnlichen Lebesgueschen Sinne meßbare Funktion ist, mit Werten > 0 , die Größe

$$\mu^*(N) = \int \dots \int_{\{x_1, \dots, x_k\} \text{ in } N} f(x_1, \dots, x_k) dx_1 \dots dx_k^{11};$$
 usw. Der allereinfachste

Fall besteht, wenn Ω ein Intervall $0 \leq x < a$ ist (a positiv-endlich oder ∞), und μ^* das Lebesguesche äußere Maß ist.

Ist ein m -Raum Ω mit einem l -Maß gegeben, so betrachten wir seine meßbaren Teilmengen. Zwei solche, M, N , nennen wir äquivalent, in Zeichen $M \sim N$, wenn sie sich nur um eine Menge vom Maße 0 unterscheiden, d. h. wenn $\mu^*((M+N) - M \cdot N) = 0$ ist. Da offenbar $M \sim M$ gilt, aus $M \sim N$ $N \sim M$ folgt, und aus $M \sim N, N \sim A$ $M \sim A$ folgt, hat dieser Begriff die wesentlichen Eigenschaften eines Äquivalenzbegriffes. Wir definieren nun:

DEFINITION 4. Seien Ω, Ω' zwei m -Räume mit den bzw. l -Mäßen μ^*, μ'^* . Eine maßtreue Punktabbildung von Ω auf Ω' ist eine eindeutige Abbildung $x' = \varphi(x)$ von Ω auf Ω' , die jeder meßbaren Teilmenge M von Ω ein meßbares Bild M' in Ω' vom gleichen Maße zuordnet, und jeder meßbaren Teilmenge M' von Ω' ein meßbares Urbild M in Ω vom gleichen Maße.

DEFINITION 5. Seien $\Omega, \Omega', \mu^*, \mu'^*$ wie vorhin. Eine maßtreue Mengenabbildung von Ω auf Ω' ist eine solche, die jeder meßbaren Teilmenge von Ω eine meßbare Teilmenge von Ω' vom gleichen Maße zuordnet, und zwar mit den folgenden Eigenschaften:

1. Zu jedem meßbaren M' in Ω' existiert ein meßbares M in Ω , dem eine M' äquivalente Menge zugeordnet ist.
2. Sind $M_1, M_2, \dots$ (in Ω) die bzw. Mengen $M'_1, M'_2, \dots$ (in Ω') zugeordnet, so sind $M_1 + M_2 + \dots, M_1 \cdot M_2 \cdot \dots$ bzw. $M'_1 + M'_2 + \dots, M'_1 \cdot M'_2 \cdot \dots$ zugeordnet. Sind M, N (in Ω) bzw. M', N' (in Ω') zugeordnet, so ist $M - N$ $M' - N'$ zugeordnet¹².

¹¹ Falls N nach Lebesgue nicht meßbar ist, bedeute $\int \dots \int_{\{x_1, \dots, x_k\} \text{ in } N}$ das „obere“ Lebesguesche Integral. $f(x_1, \dots, x_k) \equiv 1$ führt zum gewöhnlichen Lebesgueschen Maß zurück.

¹² Von diesen drei Forderungen genügen zwei, da aus ihnen die dritte folgt.

Zwei maßtreue Mengenabbildungen (von Ω auf Ω') heißen äquivalent, wenn für jedes meßbare M in Ω die durch die beiden Mengenabbildungen ihm zugeordneten Mengen in Ω' einander äquivalent sind.

Eine maßtreue Punktabbildung $x' = \varphi(x)$ erzeugt offenbar eine maßtreue Mengenabbildung: indem wir jedem meßbaren M in Ω sein durch $x' = \varphi(x)$ vermitteltes Bild M' in Ω' zuordnen¹³. Wir werden aber auch die Umkehrung hiervon beweisen: (Satz 1) zu jeder maßtreuen Mengenabbildung kann eine maßtreue Punktabbildung gefunden werden, die eine ihr äquivalente maßtreue Mengenabbildung erzeugt.

Diese, und noch einige weitere Sätze bilden den Inhalt dieser Arbeit.

2. Seien Ω, Ω' m -Räume mit den bzw. l -Maßen μ^*, μ'^* , und sei eine maßtreue Mengenabbildung von Ω auf Ω' (im Sinne von Def. 5) gegeben. Ehe wir den angekündigten Satz 1 beweisen, leiten wir den folgenden Hilfssatz ab:

HILFSSATZ. Sei M eine meßbare Menge in Ω , und es sei ein $\varepsilon > 0$ gegeben. Dann kann M in paarweise elementfremde Mengen $\Gamma_1, \Gamma_2, \dots, N$ zerlegt werden, $M = \Gamma_1 + \Gamma_2 + \dots + N$ (die Anzahl der Addenden Γ_n darf auch endlich, oder sogar 0, sein), derart, daß $\mu(N) = 0$ ist, und jedes Γ_n abgeschlossen ist, $\mu(\Gamma_n) > 0$, und einen Diameter $\leq \varepsilon$ hat¹⁴.

Die Bedingung über den Diameter ist offenbar gegenstandslos, wenn M selbst einen Diameter $\leq \varepsilon$ hat. Ist ferner $M = M_1 + M_2 + \dots$, wo $M_1, M_2, \dots$ paarweise elementfremde meßbare Mengen sind, für die der Hilfssatz gesichert ist, so steht er auch für M fest: denn

$$M_1 = \Gamma_{11} + \Gamma_{12} + \dots + N_1, \quad M_2 = \Gamma_{21} + \Gamma_{22} + \dots + N_2, \dots$$

ergibt für M die Zerlegung

$$M = \Gamma_1 + \Gamma_2 + \dots + N,$$

wo $\Gamma_1, \Gamma_2, \dots$ die $\Gamma_{11}, \Gamma_{12}, \dots, \Gamma_{21}, \Gamma_{22}, \dots, \dots$ in irgendeiner Reihenfolge sind, und $N = N_1 + N_2 + \dots$. Es genügt also, den Hilfssatz ohne die Diameterbedingung zu beweisen, wenn jedes meßbare $M = M_1 + M_2 + \dots$, mit paarweise elementfremden, meßbaren $M_1, M_2, \dots$, mit Diametern $\leq \varepsilon$ ist. Nun ist Ω separabel, sei $x_1, x_2, \dots$ eine überall dichte Folge, $K_1, K_2, \dots$ die Kugeln vom Radius ε um diese bez. Mittelpunkte. Dann ist

$$\Omega = K_1 + K_2 + \dots,$$

so daß

$$M_1 = K_1 \cdot M, \quad M_2 = (K_2 - K_1 \cdot K_2) \cdot M,$$

$$M_3 = (K_3 - (K_1 + K_2) \cdot K_3) \cdot M, \dots$$

das Gewünschte leisten.

¹³ 1. gilt sogar in verschärfter Form: jedes meßbare M' in Ω' ist selbst einem meßbaren M in Ω (seinem durch $x' = \varphi(x)$ vermittelten Urbild) zugeordnet.

¹⁴ Der Diameter einer Menge Γ ist die obere Grenze aller Zahlen $\overline{xy}$, x in Γ , y in Γ .

Wir können also die Diameterbedingung fortlassen, und M in einer Kugel K gelegen annehmen. K ist abgeschlossen, also meßbar, also $K - M$ auch. Sei $\delta > 0$, nach Def. 2, V existiert ein offenes $O \supset K - M$, mit $\mu(O) \leq \mu(K - M) + \delta$. Also ist erst recht

$$\mu(O \cdot K) \leq \mu(K - M) + \delta,$$

und

$$\mu(K - O \cdot K) = \mu(K) - \mu(O \cdot K) \geq \mu(K) - \mu(K - M) - \delta = \mu(M) - \delta.$$

Dabei ist $\Gamma = K - O \cdot K$ abgeschlossen (K ist abgeschlossen, O offen) und $\subset M$.

Bilden wir nun das Γ für M und $\delta = \frac{1}{2}$, es heiße Γ_1 , dann das Γ für $M - \Gamma_1$ und $\delta = \frac{1}{2}$, es heiße Γ_2 , dann das Γ für $M - \Gamma_1 - \Gamma_2$ und $\delta = \frac{1}{2}$, es heiße $\Gamma_3, \dots$. Nach Konstruktion ist

$$\mu(M - \Gamma_1 - \dots - \Gamma_n) \leq \frac{1}{n+1},$$

also auch

$$\mu(M - \Gamma_1 - \Gamma_2 - \dots) \leq \frac{1}{n+1}, \quad \text{d. h.} = 0.$$

Diese $\Gamma_1, \Gamma_2, \dots$ und $N = M - \Gamma_1 - \Gamma_2 - \dots$ leisten daher alles Gewünschte, wenn wir noch diejenigen unter ihnen, die das Maß 0 haben, falls solche überhaupt vorliegen, fortlassen und zu N hinzufügen. Damit ist der Hilfssatz bewiesen.

Nunmehr gehen wir so vor: Wir zerlegen Ω nach dem Hilfssatz in $\Gamma_1 + \Gamma_2 + \dots + N$ mit $\varepsilon = \frac{1}{2}$. Die Mengenabbildung ordne $\Gamma_1, \Gamma_2, \dots, N$ die Mengen $\Gamma'_1, \Gamma'_2, \dots, N'$ in Ω' zu. Je zwei der ersteren Mengen haben leere Durchschnitte, also haben je zwei der letzteren Durchschnitte vom Maße 0. Daher sind $\Gamma'_1, \Gamma'_2, \Gamma'_3, \dots$ bzw. den Mengen

$$\bar{\Gamma}'_1 = \Gamma'_1, \quad \bar{\Gamma}'_2 = \Gamma'_2 - \Gamma'_1 \cdot \Gamma'_2, \quad \bar{\Gamma}'_3 = \Gamma'_3 - (\Gamma'_1 + \Gamma'_2) \cdot \Gamma'_3, \dots$$

äquivalent. Also: $\Gamma_1, \Gamma_2, \dots$ sind paarweise elementfremd und $\subset \Omega$, ebenso $\bar{\Gamma}'_1, \bar{\Gamma}'_2, \dots$ und $\subset \Omega'$; $\Gamma_1 + \Gamma_2 + \dots$ ist $\subset \Omega$ und ihm äquivalent, $\bar{\Gamma}'_1 + \bar{\Gamma}'_2 + \dots$ ist $\subset \Omega'$ und ihm äquivalent; die Γ_n zugeordnete Menge ist $\bar{\Gamma}'_n$ äquivalent; Γ_n ist abgeschlossen, vom Maße > 0 , und sein Diameter $\leq \frac{1}{2}$.

Jetzt zerlegen wir jedes $\bar{\Gamma}'_n$ nach dem Hilfssatz in $\Gamma'_{n1} + \Gamma'_{n2} + \dots + N'_n$ mit $\varepsilon = \frac{1}{2}$. Die Mengenabbildung ordne $\Gamma'_{n1}, \Gamma'_{n2}, \dots, N'_n$ die Mengen $\Gamma_{n1}, \Gamma_{n2}, \dots, N_n$ in Ω zu. Da je zwei der ersteren Mengen leere Durchschnitte haben und $\subset \bar{\Gamma}'_n$ sind, also $\subset$ in einer Menge $\sim \Gamma'_n$, können wir die letzteren wiederum durch äquivalente Mengen ersetzen, die auch paarweise elementfremd und $\subset \Gamma_n$ sind. So werde aus $\Gamma_{n1}, \Gamma_{n2}, \dots$ bzw. $\bar{\Gamma}_{n1}, \bar{\Gamma}_{n2}, \dots$. Also: $\bar{\Gamma}_{n1}, \bar{\Gamma}_{n2}, \dots$ sind paarweise elementfremd und $\subset \Gamma_n$, ebenso $\Gamma'_{n1},$

$\Gamma'_{n2}, \dots$ und $\subset \bar{\Gamma}'_n$; $\bar{\Gamma}_{n1} + \bar{\Gamma}_{n2} + \dots$ ist $\subset \Gamma_n$ und ihm äquivalent, $\Gamma'_{n1} + \Gamma'_{n2} + \dots$ ist $\subset \bar{\Gamma}'_n$ und ihm äquivalent; die $\bar{\Gamma}_{np}$ zugeordnete Menge ist $\bar{\Gamma}'_{np}$ äquivalent; Γ'_{np} ist abgeschlossen, vom Maße > 0 , und sein Diameter $\leq \frac{1}{2}$.

Die mit Ω , Ω' und $\frac{1}{2}$ vorgenommene Konstruktion wiederholen wir an $\bar{\Gamma}_{np}$, Γ'_{np} und $\frac{1}{2}$. So entstehen Mengen Γ_{npq} , $\bar{\Gamma}'_{npq}$ und $\bar{\Gamma}_{npqr}$, Γ'_{npqr} mit den folgenden Eigenschaften: Γ_{np1} , Γ_{np2} , $\dots$ sind paarweise elementfremd und $\subset \bar{\Gamma}_{np}$, ebenso $\bar{\Gamma}'_{np1}$, $\bar{\Gamma}'_{np2}$, $\dots$ und $\subset \bar{\Gamma}'_{np}$; $\Gamma_{np1} + \Gamma_{np2} + \dots$ ist $\subset \bar{\Gamma}_{np}$ und ihm äquivalent, $\bar{\Gamma}'_{np1}$, $\bar{\Gamma}'_{np2}$, $\dots$ ist $\subset \bar{\Gamma}'_{np}$ und ihm äquivalent; die Γ_{npq} zugeordnete Menge ist $\bar{\Gamma}'_{npq}$ äquivalent; Γ_{npq} ist abgeschlossen, vom Maße > 0 , und sein Diameter $\leq \frac{1}{2}$. $\bar{\Gamma}_{npq1}$, $\bar{\Gamma}_{npq2}$, $\dots$ sind paarweise elementfremd und $\subset \Gamma_{npq}$, ebenso Γ'_{npq1} , Γ'_{npq2} , $\dots$ und $\subset \bar{\Gamma}'_{npq}$; $\bar{\Gamma}_{npq1} + \bar{\Gamma}_{npq2} + \dots$ ist $\subset \Gamma_{npq}$ und ihm äquivalent, $\Gamma'_{npq1} + \Gamma'_{npq2} + \dots$ ist $\subset \bar{\Gamma}'_{npq}$ und ihm äquivalent; die $\bar{\Gamma}_{npqr}$ zugeordnete Menge ist Γ'_{npqr} äquivalent; Γ'_{npqr} ist abgeschlossen, vom Maße > 0 , und sein Diameter $\leq \frac{1}{2}$.

Dieselbe Konstruktion wiederholen wir an $\bar{\Gamma}_{npqr}$, Γ'_{npqr} und $\frac{1}{2}$, und gewinnen dadurch Mengen Γ_{npqrs} , $\bar{\Gamma}'_{npqrs}$ und $\bar{\Gamma}_{npqrst}$, Γ'_{npqrst} usw. Zusammenfassend entsteht ein System von Mengen $\Gamma_{np\dots fg}$, $\bar{\Gamma}'_{np\dots fg}$ und $\bar{\Gamma}_{np\dots fgh}$, $\Gamma'_{np\dots fgh}$ ($n, p, \dots, f, g$ sei eine beliebige ungerade Anzahl von Indizes, etwa $2\nu - 1$, deren jeder über 1, 2, $\dots$ oder ein Anfangsstück davon läuft, $n, p, \dots, f, g, h$ entsprechend eine gerade Anzahl 2ν). Dabei gilt: $\Gamma_{np\dots f1}$, $\Gamma_{np\dots f2}$, $\dots$ paarweise elementfremd und $\subset \bar{\Gamma}_{np\dots f}$, ebenso $\bar{\Gamma}'_{np\dots f1}$, $\bar{\Gamma}'_{np\dots f2}$, $\dots$ und $\subset \bar{\Gamma}'_{np\dots f}$; $\Gamma_{np\dots f1} + \Gamma_{np\dots f2} + \dots$ ist $\subset \bar{\Gamma}_{np\dots f}$ und ist ihm äquivalent, $\bar{\Gamma}'_{np\dots f1} + \bar{\Gamma}'_{np\dots f2} + \dots$ ist $\subset \bar{\Gamma}'_{np\dots f}$ und ist ihm äquivalent; die $\Gamma_{np\dots fg}$ zugeordnete Menge ist $\bar{\Gamma}'_{np\dots fg}$ äquivalent; $\Gamma_{np\dots fg}$ ist abgeschlossen, vom Maße > 0 , und sein Diameter ist $\leq \frac{1}{\nu}$. Analoges gilt für $\bar{\Gamma}_{np\dots fgh}$, $\Gamma'_{np\dots fgh}$, nur ist dort $\Gamma'_{np\dots fgh}$ abgeschlossen, vom Maße > 0 , und vom Diameter $\leq \frac{1}{\nu}$. Für $\nu = 1$ tritt an Stelle von $\Gamma_{np\dots f}$, $\bar{\Gamma}'_{np\dots f}$ (das $2\nu - 2 = 0$ Indizes hätte) Ω , Ω' .

Alle Mengen $\bar{\Gamma}_{np\dots f} - \Gamma_{np\dots f1} - \Gamma_{np\dots f2} - \dots$, $\Gamma'_{np\dots f} - \bar{\Gamma}'_{np\dots f1} - \bar{\Gamma}'_{np\dots f2} - \dots$, $\Gamma_{np\dots fg} - \bar{\Gamma}_{np\dots fg1} - \bar{\Gamma}_{np\dots fg2} - \dots$, $\bar{\Gamma}'_{np\dots fg} - \Gamma'_{np\dots fg1} - \Gamma'_{np\dots fg2} - \dots$ haben das Maß 0, also auch ihre Vereinigungsmengen über alle ν und $n, p, \dots, f, g, h$ (der ersten und dritten in Ω , der zweiten und vierten in Ω'): M_0 , M'_0 . Aus dem bisher Gesagten folgt, daß jeder Punkt von $\Omega - M_0$ genau einem Durchschnitt $\Gamma_n \cdot \bar{\Gamma}_{np} \cdot \Gamma_{npq} \cdot \bar{\Gamma}_{npqr} \dots$ angehört, und jeder Punkt von $\Omega' - M'_0$ genau einem Durchschnitt $\bar{\Gamma}'_n \cdot \Gamma'_{np} \cdot \bar{\Gamma}'_{npq} \cdot \Gamma'_{npqr} \dots$. Andererseits ist jeder Durchschnitt $\Gamma_n \cdot \bar{\Gamma}_{np} \cdot \Gamma_{npq} \cdot \bar{\Gamma}_{npqr} \dots$ wegen $\Gamma_n \supset \bar{\Gamma}_{np} \supset \Gamma_{npq} \supset \bar{\Gamma}_{npqr} \supset \dots$ gleich $\Gamma_n \cdot \Gamma_{npq} \dots$. Da Γ_n , Γ_{npq} , $\dots$ eine absteigende Folge nicht leerer abgeschlossener Mengen ist, deren Diameter gegen 0 streben, besteht ihr Durchschnitt aus genau einem Punkt (vgl. a. a. O. Anm. ²). Ebenso zeigt

man, daß der Durchschnitt $\bar{\Gamma}'_n \cdot \Gamma'_{np} \cdot \bar{\Gamma}'_{npq} \cdot \Gamma'_{npqr} \dots$ aus genau einem Punkte besteht (hier ist $\Gamma'_{np} \cdot \Gamma'_{npqr} \dots$ heranzuziehen). Sei die Menge aller $\Gamma_n \cdot \bar{\Gamma}_{np} \cdot \Gamma_{npq} \cdot \bar{\Gamma}_{npqr} \dots$ -Durchschnittspunkte Ω_1 , die aller $\bar{\Gamma}'_n \cdot \Gamma'_{np} \cdot \bar{\Gamma}'_{npq} \cdot \Gamma'_{npqr} \dots$ -Durchschnittspunkte Ω'_1 . Da $\Omega_1 \supset \Omega - M_0$, $\Omega - \Omega_1 \subset M_0$ ist, ist $\Omega - \Omega_1$ vom Maße 0, d. h. Ω_1 dem Ω äquivalent, ebenso ist Ω'_1 dem Ω' äquivalent.

Auf diese Weise sind sowohl die Punkte von Ω_1 , als auch diejenigen von Ω'_1 , eindeutig auf die Folgen $n, p, q, r, \dots$ abgebildet; dies bewirkt auch eine eindeutige Abbildung von Ω_1 auf Ω'_1 , welche $x' = \varphi(x)$ heißen möge.

Aus der Definition dieser Abbildung folgt, daß sie $\Gamma_{np\dots fg} \cdot \Omega_1$ in $\bar{\Gamma}'_{np\dots fg} \cdot \Omega'_1$ überführt. Diese Mengen sind $\Gamma_{np\dots fg}$ bzw. $\bar{\Gamma}'_{np\dots fg}$ äquivalent, haben also bzw. dieselben Maße wie diese. $\Gamma_{np\dots fg}$ hat dasselbe Maß, wie die ihr zugeordnete Menge, diese wieder dasselbe Maß, wie das ihr äquivalente $\bar{\Gamma}'_{np\dots fg}$. Also haben $\Gamma_{np\dots fg} \cdot \Omega_1$ und $\bar{\Gamma}'_{np\dots fg} \cdot \Omega'_1$, ihr Bild, dasselbe Maß. Wenn Δ die Summe einiger Mengen $\Gamma_{np\dots fg}$ ist (alle von gleicher Indizesanzahl $2\nu - 1$, also paarweise elementfremd), so ist demnach das $x' = \varphi(x)$ -Bild von $\Delta \cdot \Omega_1$ meßbar und vom selben Maß wie diese Menge.

Sei M abgeschlossen (in Ω). $\Gamma^{(\nu)}$ sei die Summe aller $\Gamma_{np\dots fg}$ mit $2\nu - 1$ Indizes, die Punkte aus M enthalten. Es ist $\Gamma^{(\nu)} \supset M$, und da jedes $\Gamma_{np\dots fg}$ einen Durchmesser $\leq \frac{1}{\nu}$ hat, hat jeder Punkt von $\Gamma^{(\nu)}$ eine Distanz $\leq \frac{1}{\nu}$ von M . Daher ist $M = \Gamma^{(1)} \cdot \Gamma^{(2)} \dots$, ferner $\Gamma^{(1)} \supset \Gamma^{(2)} \supset \dots$; also auch

$$M \cdot \Omega_1 = (\Gamma^{(1)} \cdot \Omega_1) \cdot (\Gamma^{(2)} \cdot \Omega_1) \dots, \quad \Gamma^{(1)} \Omega_1 \supset \Gamma^{(2)} \Omega_1 \supset \dots$$

Da das Bild von $\Gamma^{(\nu)} \cdot \Omega_1$ meßbar und vom selben Maße ist, gilt dasselbe vom Bilde von $M \cdot \Omega_1$.

Wenn O offen ist, so ist $O = M_1 + M_2 + \dots$, wobei $M_1, M_2, \dots$ abgeschlossen und $M_1 \subset M_2 \subset \dots$ sind; also $O \cdot \Omega_1 = M_1 \cdot \Omega_1 + M_2 \cdot \Omega_1 + \dots$, $M_1 \cdot \Omega_1 \subset M_2 \cdot \Omega_1 \subset \dots$, und da das Bild von $M_n \cdot \Omega_1$ meßbar und vom selben Maße ist, gilt dasselbe von $O \cdot \Omega_1$. Ist $N = O_1 \cdot O_2 \dots$, wobei $O_1, O_2, \dots$ offen und $O_1 \supset O_2 \supset \dots$ sind, so gilt dasselbe wegen $N \cdot \Omega_1 = (O_1 \cdot \Omega_1) \cdot (O_2 \cdot \Omega_1) \dots$, $O_1 \cdot \Omega_1 \supset O_2 \cdot \Omega_1 \supset \dots$, auch für $N \cdot \Omega_1$. Schließlich ist $O_1 \supset O_2 \supset \dots$ überflüssig, da wir $O_1, O_2, \dots$ durch $O_1, O_1 \cdot O_2, \dots$ ersetzen können.

Wenn $\mathcal{A}$ das Maß 0 hat, so gibt es ein $M = O_1 \cdot O_2 \dots \supset \mathcal{A}$ vom Maß 0 (wegen Def. 2, V). Das Bild von $\mathcal{A} \cdot \Omega_1$ ist Teil des Bildes von $M \cdot \Omega_1$, da dieses das Maß 0 hat, hat jenes auch das Maß 0 und ist meßbar.

Die Aussage: das Bild von $M \cdot \Omega_1$ ist meßbar und hat dasselbe Maß wie M , gilt also 1. für jedes $M = O_1 \cdot O_2 \dots$, $O_1, O_2, \dots$ offen; 2. für jedes M vom Maße 0. Da jedes meßbare M gleich einem von der Art 1. minus einem von der Art 2. ist (wegen Def. 2, V, oder durch Anwendung des Hilfssatzes auf $\Omega - M$), gilt die obige Aussage für jedes meßbare M .

Also wird jedes meßbare $M \subset \Omega_1$ durch $x' = \varphi(x)$ auf ein meßbares $M' \subset \Omega'_1$ vom selben Maße abgebildet. Auf dieselbe Weise beweist man die Umkehrung. Nach Def. 1 bildet also $x' = \varphi(x)$ Ω_1 auf Ω'_1 maßtreu ab.

Für welche M ist das $x' = \varphi(x)$ -Bild von $M \cdot \Omega_1$ der zugeordneten Menge (in Ω') äquivalent? Für $M = \Gamma_{np \dots fg}$ ist das jedenfalls der Fall (das Bild ist $\bar{\Gamma}'_{np \dots fg} \cdot \Omega'_1$, äquivalent $\bar{\Gamma}'_{np \dots fg}$, also der Menge, die $\Gamma_{np \dots fg}$ zugeordnet ist); also auch für die weiter oben erwähnten Δ ; also bei abgeschlossenem M für die weiter oben erwähnten $\Gamma^{(v)}$, und für $M = \Gamma^{(1)} \cdot \Gamma^{(2)} \dots$ selbst; also, auf Grund derselben Betrachtungen wie weiter oben nacheinander 1. für jedes offene O , 2. für jedes $M = O_1 \cdot O_2 \dots$, $O_1, O_2, \dots$ offen, 3. für jedes M vom Maße 0, 4. für jedes meßbare M . Somit ist die durch $x' = \varphi(x)$ erzeugte Mengenzuordnung der ursprünglich gegebenen äquivalent.

Ein Mangel besteht noch: nämlich daß $x' = \varphi(x)$ Ω_1 auf Ω'_1 abbildet, und nicht Ω auf Ω' . Dieser wird, falls Ω, Ω' un abzählbar sind, wie folgt behoben. Sei $M_1 \subset \Omega$ vom Maße 0 und von der Mächtigkeit des Kontinuums,¹⁵ $M'_1 \subset \Omega'$ ebenso. $M'_1 \cdot \Omega'_1$ hat auch das Maß 0, also auch sein Urbild M_2 , ebenso $M_1 \cdot \Omega_1$ und sein Bild M'_2 . Somit sind $M_1 + M_2 \subset \Omega$ und $M'_1 + M'_2 \subset \Omega'$ vom Maße 0 und der Mächtigkeit des Kontinuums,¹⁶ und $(M_1 + M_2) \cdot \Omega_1$ hat das Bild $(M'_1 + M'_2) \cdot \Omega'_1$. Wir setzen

$$\Omega_2 = \Omega_1 - (M_1 + M_2) \cdot \Omega_1, \quad \Omega'_2 = \Omega'_1 - (M'_1 + M'_2) \cdot \Omega'_1,$$

dann ist Ω'_2 das Bild von Ω_2 , und da $\Omega_1 - \Omega_2 \subset M_1 + M_2$ das Maß 0 hat, hat auch $\Omega - \Omega_2$ das Maß 0, und weil es $\supset M_1 + M_2$ ist, die Mächtigkeit des Kontinuums. Dasselbe gilt für $\Omega' - \Omega'_2$.

Schränken wir nun die Abbildung $x' = \varphi(x)$ auf Ω_2, Ω'_2 ein, verwenden wir dagegen in $\Omega - \Omega_2, \Omega' - \Omega'_2$ irgendeine eineindeutige Abbildung $x' = \psi(x)$ der ersteren Menge auf die letztere (beide haben ja dieselbe Mächtigkeit). Beide Abbildungen machen zusammen eine eineindeutige Abbildung $x' = \chi(x)$ von Ω auf Ω' aus. Da sich aber $x' = \chi(x)$ von $x' = \varphi(x)$ (in Ω wie in Ω') nur auf einer Menge vom Maße 0 unterscheidet, hat die erstere

¹⁵ Eine solche Menge konstruiert man z. B. so: Nach Hausdorff, S. 138, Satz XI und S. 137, Satz X, hat das un abzählbare Ω ein „dyadisches Diskontinuum“ D zur Teilmenge. D. h. es ist jeder Folge $u_1, u_2, \dots$ ($= 0, 1$) ein Element $x_{u_1 u_2 \dots}$ von Ω zugeordnet, derart daß für jede gegebene Umgebung O von $x_{u_1 u_2 \dots}$ ein n existiert, daß alle $x_{u_1 u_2 \dots}$ mit $u_1 = v_1, \dots, u_n = v_n$ zu O gehören; und D die Menge aller $x_{u_1 u_2 \dots}$ ist. Seien nun $u_1, u_3, \dots$ fest gegeben, und $D_{u_1 u_3 \dots}$ die Menge aller $x_{u_1 u_2 \dots}$ ($u_2, u_4, \dots$ beliebig). Dabei ist $D \subset K$ für eine geeignete Kugel K . Die $D_{u_1 u_3 \dots}$ sind alle abgeschlossen, von der Mächtigkeit des Kontinuums, und $\subset K$. Je zwei solche Mengen sind elementfremd, also ist die Maßsumme irgendwelcher endlich vielen unter ihnen $\leq \mu(K)$ (endlich!). Daher können höchstens abzählbar viele unter ihnen Maße > 0 haben. Aber ihre Anzahl ist Kontinuum, also gibt es unter ihnen auch solche vom Maße 0.

¹⁶ Ω, Ω' haben selbst, als separable Mengen, höchstens die Mächtigkeit des Kontinuums.

Abbildung mit der letzteren dies gemein: sie ist maßtreu, und erzeugt eine Mengenzuordnung, welche der ursprünglich gegebenen äquivalent ist.

Damit haben wir bewiesen:

SATZ 1. Seien Ω , Ω' m -Räume mit den bzw. l -Maßen μ^* , μ'^* . Jede maßtreue Punktabbildung von Ω auf Ω' , oder auch nur von Ω_1 auf Ω'_1 mit $\Omega_1 \subset \Omega$, $\Omega'_1 \subset \Omega'$, $\Omega - \Omega_1$ und $\Omega' - \Omega'_1$ vom Maße 0, erzeugt eine maßtreue Mengenabbildung von Ω auf Ω' . Umgekehrt ist jede maßtreue Mengenabbildung von Ω auf Ω' einer solchen äquivalent, die durch eine maßtreue Punktabbildung eines Ω_1 auf ein Ω'_1 (vgl. oben) erzeugt wird. Sind Ω , Ω' un abzählbar, so können wir sogar $\Omega_1 = \Omega$, $\Omega'_1 = \Omega'$ erreichen. (Vgl. hierzu Def. 4, 5 in § 1, und das daran anschließend Gesagte.)

3. Seien Ω , Ω' , μ^* , μ'^* wie bisher, $x' = \varphi(x)$ eine eindeutige Abbildung von Ω auf Ω' , von der bloß die eine Hälfte der in Def. 4 formulierten Eigenschaften verlangt werde: das Urbild eines jeden meßbaren M' in Ω' soll ein meßbares M in Ω vom selben Maße sein. Wir wollen zeigen, daß dennoch die ganze Def. 4 (d. h. auch die umgekehrte Eigenschaft) gilt, d. h. daß $x' = \varphi(x)$ maßtreu ist.

Sei $x'_1, x'_2, \dots$ eine in Ω' überall dichte Folge, $K'_{\nu\lambda}$ ($\nu, \lambda = 1, 2, \dots$) die Kugel mit dem Mittelpunkt x'_ν und dem Radius $\frac{1}{\lambda}$. Das Urbild von $K'_{\nu\lambda}$ ist meßbar, also bis auf eine Menge $N_{\nu\lambda}$ vom Maße 0 eine Borelsche Menge (vgl. den Hilfssatz in § 2), diese heiße $A_{\nu\lambda}$. Die Summe aller $N_{\nu\lambda}$ ist auch vom Maße 0, sie ist Teil einer Borelschen Menge M_{00} vom Maße 0 (vgl. die analogen Schlüsse in § 2). In $\Omega - M_{00}$ gilt also: $\varphi(x)$ liegt dann und nur dann in $K'_{\nu\lambda}$, wenn x in $A_{\nu\lambda}$ liegt, d. h. soweit es in $A_{\nu\lambda} - M_{00}$ liegt. Für die auf $\Omega - M_{00}$ eingeschränkte Funktion $\varphi(x)$, sie heiße $\varphi_1(x)$, gilt also: das Urbild von M' (in Ω' , M' braucht nicht nur aus $\varphi_1(x)$ -Punkten zu bestehen, sein Urbild sind jene x , deren $\varphi_1(x)$ in M' liegen, — übrigens soll im Folgenden der Begriff „Urbild“ stets so verstanden werden) ist eine Borelsche Menge, falls M' ein $K'_{\nu\lambda}$ ist. Da jedes offene M' Summe einer Folge von $K'_{\nu\lambda}$'s ist, gilt dies auch für alle offenen M' , und dann auch für deren Komplementären, die abgeschlossenen M' . Hieraus wieder schließt man, daß es für alle Borelschen M' gilt.

Sei Ω'' der in § 1 als drittes Beispiel eines m -Raumes erwähnte ∞ -dimensionale Raum aller $\{\xi_1, \xi_2, \dots\}$ (vgl. auch Anm. ⁴). In dem wir jedem x' von Ω' das $\psi(x') = \{x'_1, x', x'_2, x', \dots\}$ von Ω'' zuordnen, bilden wir Ω' eindeutig und mitsamt seiner Umkehrung stetig auf eine Teilmenge Ω''_0 von Ω'' ab. In der Tat: Aus $\psi(x'^{(1)}) = \psi(x'^{(2)})$ folgt $\overline{x'_n, x'^{(1)}} = \overline{x'_n, x'^{(2)}}$, wenn wir also eine Folge $x'_n, x'_n, \dots$ wählen, die gegen $x'^{(1)}$ strebt, so ist $\overline{x'_n, x'^{(1)}} \rightarrow 0$, $\overline{x'_n, x'^{(2)}} \rightarrow 0$, d. h. $x'_n, x'_n, \dots$ strebt

auch gegen $x^{(2)}, x^{(1)} = x^{(2)}$. Aus $x^{(m)} \rightarrow x'$ folgt $\overline{x'_n, x^{(m)}} \rightarrow \overline{x'_n, x}$ (für $m \rightarrow \infty$), also $\psi(x^{(m)}) \rightarrow \psi(x')$ (vgl. Anm. 4). Aus $\psi(x^{(m)}) \rightarrow \psi(x')$, d. h. $\overline{x'_n, x^{(m)}} \rightarrow \overline{x'_n, x'}$ für jedes n folgt: sei $\delta > 0$, x'_n mit $\overline{x'_n, x'} \leq \delta$ gewählt, dann ist für ein genügend großes m $\overline{x'_n, x^{(m)}} \leq 2\delta$, also $\overline{x', x^{(m)}} \leq \delta + 2\delta = 3\delta$ — somit gilt $x^{(m)} \rightarrow x'$.

Als eindeutiges stetiges Bild des m -Raumes Ω' ist Ω'_0 (absolut) Borelsch¹⁷. Sei $M'' \subset \Omega''$ (absolut) Borelsch, dann ist es auch $M'' \cdot \Omega'_0$, das ψ -Urbild von M'' ist das durch die Inverse von ψ vermittelte Bild von $M'' \cdot \Omega'_0$, also wieder (absolut) Borelsch¹⁸. Das $\psi(\varphi_1)$ -Urbild von M'' ist somit das φ_1 -Urbild einer Borelschen Menge in Ω' , also, wie wir w. o. zeigten, selbst eine Borelsche Menge in Ω . Somit bildet $\psi(\varphi_1(x)) \Omega - M_{00}$ eindeutig auf eine gewisse Teilmenge von Ω'' ab, und jede Borelsche Teilmenge von Ω'' hat als Urbild eine Borelsche Menge in Ω . Daher ist $\psi(\varphi_1(x))$ eine Bairesche Funktion¹⁸. Infolgedessen hat jedes Borelsche $M \subset \Omega - M_{00}$ ein $\psi(\varphi_1)$ -Bild, das ebenfalls Borelsch ist¹⁹, und das ψ -Urbild desselben ist, wie wir weiter oben zeigten, ebenfalls Borelsch. Aber dies ist das φ_1 -Bild von M , das mit dessen φ -Bild zusammenfällt, diese sind mithin als Borelsch erwiesen.

Insbesondere ist das Bild des ganzen $\Omega - M_{00}$, es heiße $\Omega' - M'_{00}$, Borelsch, daher ist auch M'_{00} Borelsch und infolgedessen meßbar²⁰. Sein Maß ist dasselbe wie dasjenige seines φ -Urbildes M_{00} , also 0. Das φ -Bild M' eines Borelschen M ist die Summe derer von $M - M_{00} \cdot M$ und von $M_{00} \cdot M$. Ersteres ist ein φ_1 -Bild, also Borelsch, also meßbar; letzteres $\subset M'_{00}$, also ebenso wie M'_{00} vom Maße 0, also auch meßbar. Daher ist auch M' meßbar, sein Maß ist gleich demjenigen seines Urbildes M .

Ist M vom Maße 0, so ist es $\subset N$, wobei N Borelsch und vom Maße 0 ist (vgl. die entsprechenden Betrachtungen in § 2); das φ -Bild M' von M ist Teil des φ -Bildes N' von N , also wie jenes vom Maße 0, und daher meßbar. Ist schließlich M nur als meßbar vorausgesetzt, so ist es als Summe einer Borelschen Menge und einer vom Maße 0 darstellbar (vgl. den Hilfs-

¹⁷ Vgl. Hausdorff, S. 208—209, insbesondere Satz II.

¹⁸ Vgl. Hausdorff (a. a. O. Anm. 2), S. 260, Satz V. Allerdings müßte $f(x) = \psi(\varphi_1(x))$ reelle Zahlen als Werte haben, obwohl es in Wahrheit ein $\{\xi_1, \xi_2, \dots\}$, d. h. $\{\xi_1(x), \xi_2(x), \dots\}$, ist. Indessen ist der genannte Satz auf jedes $\xi_n(x)$ direkt anwendbar, und überträgt sich von diesen auf $f(x)$ (vgl. Anm. 4). Mit $\varphi_1(x)$ selbst wäre dieser Schluß nicht möglich gewesen, da es Werte aus Ω' annimmt, das in keiner Beziehung zu den reellen Zahlen steht. Vgl. hierzu a. a. O., S. 268 unten, 269 oben.

¹⁹ Vgl. Hausdorff (a. a. O. Anm. 2), S. 269, Satz XIII.

²⁰ Dies ist die Hauptschwierigkeit des Beweises: denn wohl mußte das Urbild N eines meßbaren N' meßbar sein, wir konnten aber aus der Meßbarkeit von N nicht direkt auf diejenige des Bildes N' schließen.

satz in § 2), sein φ -Bild M' ist die Summe der φ -Bilder jener Mengen, also auch meßbar. Das Maß von M' ist jenem seines Urbildes M gleich.

Damit ist die Maßtreue der Abbildung $x' = \varphi(x)$ bewiesen. Wären wir von der anderen Hälfte von Def. 4 ausgegangen, wonach das Bild jedes meßbaren M in Ω ein meßbares M' in Ω' vom selben Maße ist, so wäre dies für die Inverse von φ unsere frühere Prämisse gewesen. Daher hätte sich die Maßtreue dieser Inversen, also auch diejenige von φ , ergeben. Wir haben also bewiesen:

SATZ 2. Seien $\Omega, \Omega', \mu^*, \mu'^*$ wie bisher, $x' = \varphi(x)$ eine eindeutige Abbildung von Ω auf Ω' . Von den zwei Bedingungen in Def. 4 (Schluß von $M \subset \Omega$ auf $M' \subset \Omega'$, Schluß von $M' \subset \Omega'$ auf $M \subset \Omega$) folgt jede aus der anderen; d. h. jede von ihnen ist auch allein für die Maßtreue von φ notwendig und hinreichend.

4. Zum Schluß soll ein Satz aus einer früheren Abhandlung des Verfassers²¹ verallgemeinert und verschärft werden. Dazu muß eine dortige Begriffsbildung (d. i. die Maßfunktion, vgl. a. a. O., S. 193 unten) verallgemeinert werden.

DEFINITION 6. Sei Ω ein m -Raum, μ^* ein l -Maß in ihm, u. zw. derart, daß $\mu(\Omega)$ (also jedes $\mu^*(M)$) endlich ist. Eine Funktion $f(x)$, die in Ω definiert ist und reelle Zahlen als Werte hat, heiße meßbar, wenn jedes Intervall $\xi \leq a$ eine meßbare x -Menge zur $\xi = f(x)$ -Urbildmenge hat²². Die ganze Theorie des Lebesgueschen Integrals kann offenbar, unter Zugrundelegung des Maßes μ^* (bzw. μ), auf diese meßbaren Funktionen f übertragen werden. Das so entstehende Integral bezeichnen wir mit $\int_{\Omega} f(x) d\mu_x$.

Ist $f(x)$ meßbar, nie negativ, und beschränkt, so definieren wir für alle a , $0 \leq a \leq \text{Max}(f(x))$, eine Funktion $\varphi(a) = \mu(\text{Menge der } x \text{ mit } f(x) \leq a)$. $\varphi(a)$ ist monoton nichtfallend und nach rechts halbstetig. Ferner ist

$$0 \leq \varphi(a) \leq \text{Max}(\varphi(a)) = \mu(\Omega).$$

Da auch $\varphi(a)$ meßbar ist, und in seinem Definitionsbereiche, dem m -Raume, $0 \leq a \leq \text{Max}(f(x))$ das gewöhnliche Lebesguesche Maß als l -Maß angesehen werden kann, können wir auch seine Maßfunktion $\psi(b)$ bilden. $\psi(b)$ ist (vgl. w. o.) in $0 \leq b \leq \mu(\Omega)$ definiert, monoton nichtfallend und nach rechts halbstetig, und es ist stets $0 \leq \psi(b) \leq \text{Max}(f(x))$. Übrigens ist $\psi(b)$ die obere Grenze aller a mit $\varphi(a) \leq b$, also in einem gewissen Sinne die

²¹ Annals of Math., 3212 (1931), S. 194, Satz 1.

²² Offenbar gilt dies dann auch für die Summe aller $\xi \leq b - \frac{1}{n}$ -Intervalle ($n = 1, 2, \dots$), d. i. $\xi < b$, und für die Differenzmenge $a < \xi < b$. Hieraus folgert man es mühelos nacheinander für alle offenen, abgeschlossenen, und Borelschen ξ -Mengen.

Inverse von $\varphi(a)$; dasselbe gilt, wenn $\varphi(a)$ und $\psi(b)$ vertauscht werden (vgl. a. a. O., S. 193 unten).

Der zu beweisende Satz aber lautet so:

SATZ 3. Seien Ω, μ^* wie in Def. 6; f, φ, ψ ebenfalls wie dort, d. h. $f(x)$ in Ω definiert und mit reellen Zahlen als Werten, nicht negativ und beschränkt, $\varphi(a)$ seine Maßfunktion, $\psi(b)$ diejenige von $\varphi(a)$.

Sei $g(u)$ für reelle Zahlen definiert und mit solchen als Werten, aber sonst ganz beliebig²³. $g(f(x))$ ist dann und nur dann meßbar (vgl. Def. 6), wenn $g(\psi(b))$ es ist (dieses im gewöhnlichen Lebesgueschen Sinne), und es gilt dann

$$\int_{\Omega} g(f(x)) dv_x = \int_0^{\mu(\Omega)} g(\psi(b)) db$$

(d. h. wenn eine Seite Sinn hat, dann auch die andere, und sie sind einander gleich).

Dieser Satz ist in zwei Beziehungen allgemeiner, als der in Anm. ²¹ angeführte: Erstens braucht Ω nicht, wie dort, eine Zahlenmenge zu sein, und μ^* das gewöhnliche Lebesguesche äußere Maß²⁴; zweitens waren dort alle Schlußweisen nur in einer Richtung (von $g(\psi(b))$ auf $g(f(x))$) begründet, während die in der anderen Richtung fehlten²⁵. Wir gehen nun zu seinem Beweise über.

Es genügt offenbar, zu zeigen: wenn von den zwei Mengen $g(f(x)) \leq c$ und $g(\psi(b)) \leq c$ (c fest) die eine meßbar ist, so ist es auch die andere, und beide haben dasselbe Maß. Oder wenn Ξ die u -Menge $g(u) \leq c$ ist: wenn vom $u = f(x)$ -Urbild und vom $u = \psi(b)$ -Urbild von Ξ das eine meßbar ist, so ist es auch das andere, und beide haben dasselbe Maß. Wir beweisen dies für alle Ξ .

Betrachten wir das Zahlenintervall $0 \leq u \leq \text{Max}(f(x))$. Es ist ein m -Raum, Ω_0 . Wir definieren in ihm zwei l -Maße: erstens μ_1^* durch $\mu_1^*(\Xi) = \mu^*$ des f -Urbildes von Ξ , zweitens μ_2^* durch $\mu_2^*(\Xi) =$ gewöhnliches Lebesguesches äußeres Maß des ψ -Urbildes von Ξ . Die zu beweisende Behauptung lautet dann: μ_1^* ist mit μ_2^* identisch, oder: $\varphi_0(x) = x$ ist eine maßtreue Abbildung des Ω_0 mit μ_1^* auf das Ω_0 mit μ_2^* (vgl. Def. 4). Nach Satz 2 genügt es hierzu, zu zeigen, daß die zweite Bedingung von

²³ Es braucht z. B. keineswegs meßbar zu sein!

²⁴ Bei den Anwendungen zum Beweise des Ergodensatzes (vgl. den nächstfolgenden Artikel des Verfassers sowie den in Proc. Nat. Ac. 18, January 1931) ist gerade diese Allgemeinheit wesentlich. Am zweit-a. O. wurde dies, mittels der maßtreuen Abbildbarkeit der auftretenden Mengen auf reelle Zahlenmengen umgangen. Nun kann zwar das Bestehen dieser Abbildbarkeit, nach einer mündlichen Mitteilung von Herrn M. H. Stone an den Verfasser, allgemein bewiesen werden — es ist aber wohl von Interesse, auch einen allgemeinen direkten Beweis zu besitzen, besonders da er ebenso verläuft wie der Beweis des Spezialfalles.

²⁵ Dies reicht zwar für die Anwendungen aus, ist aber prinzipiell unbefriedigend.

Def. 4 erfüllt ist, welche im vorliegenden Falle dies besagt: wenn Ξ μ_2^* -meßbar ist, so ist es auch μ_1^* -meßbar, und die beiden Maße sind einander gleich. Für welche Ξ trifft dies nun wirklich zu?

Wenn Ξ das Intervall $u \leq a$ ist (d. h. der in $0 \leq u \leq \text{Max}(f(x))$ gelegene Teil desselben), so ist das f -Urbild die Menge $f(x) \leq a$, also meßbar, und vom Maße $\varphi(a)$, daher ist Ξ μ_1^* -meßbar und $\mu_1^*(\Xi) = \varphi(a)$. Das ψ -Urbild dagegen ist die Menge $\psi(b) \leq a$, d. h. $b <$ oder $\leq \varphi(a)$ (natürlich $b \geq 0$), also meßbar und vom Maße $\varphi(a)$, daher ist Ξ auch μ_2^* -meßbar, und $\mu_2^*(\Xi) = \varphi(a)$. Das Intervall $u \leq a$ ist also ein solches Ξ .

Also auch die Summe der Intervalle $u \leq b - \frac{1}{n}$ ($n = 1, 2, \dots$), d. i. $u < b$; also auch die Differenzmenge $a < u < b$ (vgl. Anm.²²). Und hieraus folgt dasselbe nacheinander für alle offenen, abgeschlossenen und Borelschen u -Mengen (vgl. Anm.²²). Ist $\mu_2^*(\Xi) = 0$, so ist, wie wir schon oft schlossen, $\Xi \subset H$, H Borelsch, $\mu_2(H) = 0$. Nach dem soeben Gesagten ist also auch $\mu_1(H) = 0$, also auch Ξ μ_1 -meßbar und $\mu_1(\Xi) = 0$. Also: auch die Ξ mit $\mu_2^*(\Xi) = 0$ haben die obige Eigenschaft. Und da, wie wir schon oft schlossen, jedes μ_2^* -meßbare Ξ Summe eines Borelschen und eines mit dem μ_2^* -Maß 0 ist, haben es auch diese. D. h.: die genannte Eigenschaft liegt bei allen Ξ , die für sie überhaupt in Frage kommen, vor. Damit ist der Beweis durchgeführt.

ZUR OPERATORENMETHODE IN DER KLASSISCHEN MECHANIK

Einleitung.

1. Die von Herrn B. O. Koopman entdeckte Operatorenbehandlung der klassischen Mechanik² hat in der kurzen Zeit seit ihrem Bekanntwerden eine sehr beachtenswerte Durchschlagskraft an bis dahin unbezwungenen Problemen dieser Disziplin bewiesen³. Die Hoffnung, daß sie in der Zukunft Wesentliches zum Fortschritt derselben beitragen wird, erscheint daher als berechtigt. Ihr weiterer Ausbau, der in erster Linie eine nähere Analyse der durch sie eingeführten Operatoren erheischt, ist darum höchst erwünscht. Einen Ansatz zu diesem Programm stellt die vorliegende Abhandlung dar, indem sie die hauptsächlichsten Zusammenhänge zwischen mechanischen Fragen und operatorenthoretischen Problemen angibt und einige der letzteren löst, für andere aber die gegenwärtig wahrscheinlichsten Vermutungen aufstellt.

Wir verwenden die Bezeichnungen von Koopman (vgl. a. a. O. Anm. ², ³): φ sei der Phasenraum eines mechanischen Systems, P sein allgemeiner Punkt, $S_t: P \rightarrow P_t = S_t P$ die durch die den mechanischen Differentialgleichungen gemäß in der Zeit t erfolgte Bewegung beschriebene „Strömung“ in φ . Nach dem Liouvilleschen Satze ist S_t maßtreu: d. h. jede (im Lebesgueschen Sinne) meßbare Teilmenge von φ wird durch S_t auf eine ebensolche gleichen Maßes abgebildet⁴ und umgekehrt⁵. Sei Ω eine gegenüber allen S_t invariante Teilmenge von φ , d. h. eine „Integralfläche“ der mechanischen Bewegung, man führt dann auf dieser Fläche Ω in bekannter

² Proc. Nat. Ac. 17, Mai 1931, S. 315.

³ Vgl. die Abhandlung des Verfassers Proc. Nat. Ac. 18, Jan. 1932, S. 70 (Beweis des Ergodensatzes) sowie B. O. Koopman und der Verfasser, Proc. Nat. Ac. 18, 1932, S. (Beweis des „Verrührungssatzes“).

⁴ In der Mechanik wird dies in der Regel nur für Parallelepipede bewiesen. Von diesen kann man es aber auf die Summen endlich vieler Parallelepipede übertragen sowie auf die Limites aufsteigender Folgen solcher Mengen, d. i. auf alle offenen Mengen. Sind $O_1, O_2, \dots$ offen, so gilt es auch für die Durchschnitte $O_1, O_1 O_2, O_1 O_2 O_3, \dots$ (diese sind offen), also auch für ihren Limes $O_1 O_2 \dots$. Da jede Menge vom Maße 0 Teil eines $O_1 O_2 \dots$ vom Maße 0 ist, gilt es auch für diese. Und da jede meßbare Menge Differenz eines $O_1 O_2 \dots$ und einer Menge vom Maße 0 ist, gilt es allgemein.

⁵ Die Umkehrung folgt aus der entsprechenden Eigenschaft der inversen Abbildung S_{-t} , oder auch direkt nach Satz 2 der vorhergehenden Abhandlung des Verfassers.

Weise ein Volumenelement $d\omega$ und eine Dichte $\varrho = \varrho(P) > 0$ ein, so daß das Gewicht $\varrho d\omega$ S_t -invariant ist: d. h. wenn für die Teilmengen $\mathcal{A}$ von Ω ein Lebesguesches Maß durch $\mu(\mathcal{A}) = \int_{\mathcal{A}} \varrho d\omega$ definiert wird, so ist S_t auch in Ω maßtreu⁶.

Die S_t bilden in φ (wie in Ω) eine einparametrische Gruppe, d. h. es ist stets: $S_t(S_s P) = S_{t+s} P$, symbolisch: $S_t S_s = S_{t+s}$. Sie sind, wie wir sahen, eindeutige und maßtreue Abbildungen von φ (und von Ω) auf sich selbst. Außerdem sind sie in ihrer Abhängigkeit von P, t stetig, ja sogar stetig differentiierbar.

Diese letztere Eigenschaft werden wir im folgenden nie benutzen, sie ist für die Koopmansche Operatorenmethode unwesentlich. Wir wollen daher das allgemeine Problem, das über die mechanischen „Strömungen“ hinausgeht, gleich formulieren. Den Phasenraum φ , bzw. die Integralfächen Ω , verallgemeinern wir so:

D₁. Sei Ω ein beliebiger m -Raum, d. h. metrisch separabel, und vollständig⁷, für dessen Teilmengen $\mathcal{A}$ ein l -Maß $\mu^(\mathcal{A})$ (d. h. eines mit den Eigenschaften des Lebesgueschen äußeren Maßes) definiert ist⁸.*

Die reellen Zahlen t bilden ebenfalls einen m -Raum T mit dem gewöhnlichen Lebesgueschen äußeren Maß als l -Maß, $\tilde{\mu}^(\tilde{\mathcal{A}})$.*

Der „Produktraum“ $\Omega \times T$ aller Paare $[P, t]$ ⁹ ist ebenfalls ein m -Raum, und ein l -Maß kann in ihm so definiert werden:

a) Seien O in Ω und $\tilde{O}$ in T offene Mengen, $O \times \tilde{O}$ die Menge der $[P, t]$, P in O , t in $\tilde{O}$. Dann setzen wir

$$\mathfrak{J}(O \times \tilde{O}) = \mu(O) \tilde{\mu}(\tilde{O}).$$

β) Sei $\bar{\mathcal{A}}$ eine Menge in $\Omega \times T$. Betrachten wir alle Folgen $O_1, \tilde{O}_1, O_2, \tilde{O}_2, \dots$ mit $\bar{\mathcal{A}} \subseteq O_1 \times \tilde{O}_1 + O_2 \times \tilde{O}_2 + \dots$, die untere Grenze ihrer $\mathfrak{J}(O_1 \times \tilde{O}_1) + \mathfrak{J}(O_2 \times \tilde{O}_2) + \dots$ heiße $\bar{\mu}^(\bar{\mathcal{A}})$.*

Diese Bildung ist dem mehrdimensionalen Lebesgueschen äußeren Maße analog, und sein l -Maßcharakter wird entsprechend bewiesen¹⁰.

Der Sinn von Begriffen wie „meßbare Funktion“, „Integral“ usw. in $\Omega, T, \Omega \times T$ ist nunmehr klar.

⁶ Die allgemein übliche Diskussion beschränkt sich auch hier auf Parallelepipede, was dann nach Anm. ⁴, ⁵ verallgemeinert werden kann.

⁷ Vgl. Def. 1 der vorhergehenden Abhandlung des Verfassers.

⁸ Vgl. Def. 2 a. a. O. Anm. ⁷.

⁹ Vgl. z. B. Hausdorff, „Mengenlehre“, Berlin 1927, S. 102. Ist $\overline{P, Q}$ die Distanz in Ω , so ist die in $\Omega \times T$ etwa $\overline{[P, t], [Q, s]} = \sqrt{(\overline{P, Q})^2 + (t-s)^2}$.

¹⁰ Vgl. z. B. Carathéodory, „Reelle Funktionen“, Berlin 1918, S. 229–237 sowie 274 ff. Die weiteren Eigenschaften des mehrdimensionalen Maßes übertragen sich gleichfalls ohne weiteres auf unseren Fall, insbesondere der Satz von Fubini, S. 621–625.

Das Analogon der „Strömung“ S_t aber ist dies:

D_2 . Für jedes t sei eine eindeutige maßtreue Abbildung $S_t: P \rightarrow P_t = S_t P$ von Ω auf sich selbst gegeben. Die S_t mögen eine einparametrische Gruppe bilden, d. h. es sei stets: $S_t(S_s P) = S_{t+s} P$, symbolisch $S_t S_s = S_{t+s}$.

Als P , t -Funktion sei $S_t P$ meßbar¹¹.

Es erweist sich u. U. als zweckmäßig, noch etwas weniger zu verlangen:

D'_2 . Mengen vom Maße 0 (wir nennen sie $N'_t, N''_t, N_{t,s}$) werden in D_2 als Ausnahmen zugelassen: d. h. S_t braucht nur $\Omega - N'_t$ auf $\Omega - N''_t$ eindeutig und maßtreu abzubilden (in N'_t sei S_t undefiniert); $S_t(S_s P) = S_{t+s} P$ braucht nur außerhalb von $N_{t,s}$ zu gelten¹². Letzteres heiße symbolisch $S_t S_s \cong S_{t+s}$.

Eine Schar S_t nach D'_2 nennen wir eine „allgemeine Strömung“. Gilt sogar D_2 und ist $S_t P$ in P , t stetig oder stetig differenzierbar usw., so heiße die Schar S_t bzw. eine „stetige Strömung“ eine „stetig differenzierbare Strömung“ usw. Entspringt sie schließlich, wie es eingangs der Fall war, einem mechanischen Problem, so heiße sie eine „mechanische Strömung“. — Zwei allgemeine Strömungen gelten als äquivalent, wenn sich ihre S_t für jedes gegebene t nur in einer P -Menge vom Maße 0 unterscheiden.

Eine weitere wichtige Begriffsbildung ist diese:

D_3 . Ω', Ω'' seien m -Räume mit l -Maßen (μ_1^* bzw. μ_2^*) nach D_1 . $\Sigma: P' \rightarrow P'' = \Sigma P'$ sei eine eindeutige und maßtreue Abbildung von Ω' minus eine Menge vom Maße 0 auf Ω'' minus eine Menge vom Maße 0. Jeder allgemeinen Strömung S'_t in Ω' ordnet dieses Σ eine allgemeine Strömung $S''_t = \Sigma S'_t \Sigma^{-1}$ in Ω'' zu und umgekehrt.

Zwei allgemeine Strömungen S'_t, S''_t (in Ω' bzw. Ω'') sollen isomorph heißen, wenn ein solches Σ existiert, das sie (bis auf eine Menge vom Maße 0) ineinander überführt: $S''_t \cong \Sigma S'_t \Sigma^{-1}$.

Isomorphe Strömungen können wir als nicht wesentlich verschieden ansehen, da ihr einziger Unterschied darin liegt, daß die Punkte P' von Ω' durch die entsprechenden $P'' = \Sigma P'$ von Ω'' ersetzt wurden. Bei solchen Problemen aber, bei denen es nur auf die Volumina der auftretenden Mengen ankommt (wir werden weiter unten auseinandersetzen, welche Probleme wir meinen), ist dies wegen der Maßtreue von Σ unwesentlich.

Die Stetigkeit, oder gar stetige Differenzierbarkeit einer Strömung kann bei Ersetzung durch eine isomorphe verlorengehen — dies sind keine iso-

¹¹ D. h.: wenn O eine offene Menge in Ω ist, so sei die durch $S_t P$ in O bestimmte $[P, t]$ -Menge in $\Omega \times T$ meßbar.

¹² Da $N'_t, N''_t, N_{t,s}$ von t, s abhängen und diese über un abzählbar viele Werte laufen, können wir die genannten Ausnahmemengen nicht zu einer Gesamtmenge vom Maße 0 zusammenfassen.

morphieinvarianten Eigenschaften. Vermutlich kann sogar zu jeder allgemeinen Strömung eine isomorphe stetige Strömung gefunden werden¹³, vielleicht sogar eine stetig-differentiierbare, oder gar eine mechanische. Dies mag es rechtfertigen, daß hier an Stelle der eigentlich interessanten mechanischen Strömungen alle allgemeinen untersucht werden.

2. Die mechanischen Probleme, die mit der Koopmanschen Methode angreifbar sind, sind vom folgenden Typus: Welche geometrische Wahrscheinlichkeit im Phasenraume haben gewisse, mit Hilfe der mechanischen Bewegung beschriebene Ereignisse — d. h. welches Maß haben gewisse, mit Hilfe der S_t definierte Teilmengen von Ω ? Es handelt sich also um isomorphieinvariante Eigenschaften der S_t . Einige Beispiele¹⁴:

a) Sei $\mathcal{A}$ gegeben, $Z(\mathcal{A}, P, t, s) = \tilde{\mu}(\tau\text{-Menge der } S_\tau P \text{ in } \mathcal{A}, t \leq \tau < s)$.

Gilt $\frac{Z(\mathcal{A}, P, t, s)}{s - t} \rightarrow \frac{\mu(\mathcal{A})}{\mu(\Omega)}$ bei $s - t \rightarrow +\infty$ für alle P mit Ausnahme einer Menge vom Maße 0?

b) Gilt $\frac{Z(\mathcal{A}, P, t, s)}{s - t} \rightarrow \frac{\mu(\mathcal{A})}{\mu(\Omega)}$ bei $s - t \rightarrow +\infty$ wenigstens im Sinne der Mittelkonvergenz einer P -Funktion¹⁵?

c) Seien $\mathcal{A}, M$ gegeben, $Z'(\mathcal{A}, M, t, s) = \int_t^s \mu(S_t \mathcal{A} \cdot M) \cdot d\tau$. Gilt $\frac{Z'(\mathcal{A}, M, t, s)}{s - t} \rightarrow \frac{\mu(\mathcal{A})\mu(M)}{\mu(\Omega)}$ bei $s - t \rightarrow +\infty$?

d) Gilt sogar $\mu(S_t \mathcal{A} \cdot M) \rightarrow \frac{\mu(\mathcal{A})\mu(M)}{\mu(\Omega)}$ bei $t \rightarrow \pm\infty$?

e) Gilt wenigstens $\mu(S_t \mathcal{A} \cdot M) \rightarrow \frac{\mu(\mathcal{A})\mu(M)}{\mu(\Omega)}$ bei $t \rightarrow \pm\infty$, wenn eine feste t -Menge I von der Dichte 0¹⁶ ausgeschlossen wird?

a), b), c) sind verschiedene Fassungen der sog. Quasi-Ergodenhypothese¹⁷, u. zw. von abnehmender Schärfe: aus a) folgt offenbar b), aus b), wie wir

¹³ Der Verfasser hofft, hierfür demnächst einen Beweis anzugeben.

¹⁴ Mit $\mathcal{A}, M$ werden wir meßbare Teilmengen von Ω mit endlichem Maße verstehen, $S_t \mathcal{A}$ ist das durch S_t vermittelte Bild von $\mathcal{A}$, $\mathcal{A} \cdot M$ der Durchschnitt von $\mathcal{A}$ und M .

¹⁵ D. h. ist $\int_{\Omega} \left(\frac{Z(\mathcal{A}, P, t, s)}{s - t} - \frac{\mu(\mathcal{A})}{\mu(\Omega)} \right)^2 d\nu_P \rightarrow 0$ für $s - t \rightarrow +\infty$. Mit $\int_{\Omega} \dots d\nu_P$ bezeichnen wir das auf unser μ gegründete Integral. Natürlich ist auch dies eine Aussage über Maße.

¹⁶ Dichte 0 bedeutet: $\frac{\tilde{\mu}(\tau\text{-Menge der } \tau \text{ in } I, -t \leq \tau \leq t)}{2t} \rightarrow 0$ für $t \rightarrow +\infty$. In den Fällen, die uns interessieren, ist I als Vereinigung einer sich im Unendlichen häufenden Intervallfolge wählbar.

¹⁷ Dieselbe ist in der physikalischen Literatur, trotz ihrer Wichtigkeit, mathematisch recht ungenau formuliert. Vgl. dazu die Abhandlungen des Verfassers Proc. Nat. Ac. 18, Jan. 1932, S. 70 und 18, March 1932, S. 263.

sehen werden, c). Für die physikalischen Anwendungen der Quasi-Ergodenhypothese¹⁸ ist, wie eine nähere Betrachtung lehrt, b) die notwendige und hinreichende Fassung¹⁸. d), e) sind Verschärfungen von c), also Verschärfungen der Quasi-Ergodenhypothese (für die Fassungen a), b) kommt offenbar etwas Derartiges nicht in Frage), u. zw. folgt e) aus d). Ihr Sinn ist dieser: jede Teilmenge $\mathcal{A}$ von Ω wird durch die Strömung S_t derartig verschmiert, daß sie für große t in jedes M eindringt, und Ω homogen erfüllt $\left(\frac{\mu(S_t \mathcal{A} \cdot M)}{\mu(\mathcal{A})} \rightarrow \frac{\mu(M)}{\mu(\Omega)} \right)$, d. h. in M eingedrungener relativer

Anteil von $\mathcal{A} \rightarrow$ relativer Anteil von M an Ω); oder: die Wahrscheinlichkeit dafür, daß P zur Zeit $\tau = 0$ in $\mathcal{A}$ und zur Zeit $\tau = t$ in M liegt $\left(\frac{\mu(S_t \mathcal{A} \cdot M)}{\mu(\Omega)} \right)$, P bewegt sich mit der Strömung S_t strebt gegen die Wahrscheinlichkeit dessen, daß zwei unabhängige Punkte P, Q in $\mathcal{A}$ bzw. M liegen $\left(\frac{\mu(\mathcal{A}) \mu(M)}{(\mu(\Omega))^2} \right)$; d. h.: das Wahrscheinlichkeitsverhalten von $P, S_t P$ strebt gegen die statistische Unabhängigkeit.

Die notwendigen und hinreichenden Bedingungen für die Gültigkeit von b), von c) (es sind dieselben), sowie von e), und einiges über d), konnten mit Hilfe der Koopmanschen Methode aufgefunden werden¹⁹.

Eine naheliegende Frage ist, warum wir in a)—e) alle meßbaren $\mathcal{A}, M$ zulassen, und nicht nur solche, die bei Beantwortung physikalischer Fragen tatsächlich in Frage kommen, d. h. für welche die Zugehörigkeit von P durch Messungen beschränkter Genauigkeit mit hoher Wahrscheinlichkeit entschieden werden kann. Also etwa abgeschlossene oder offene $\mathcal{A}$, oder gar nur Parallelepipede. Der Grund hierfür ist ein zweifacher:

α) Entgegen der allgemeinen Auffassung ist für jedes meßbare $\mathcal{A}$ die Zugehörigkeit von P im obigen Sinne „physikalisch“ entscheidbar.

β) Wenn eines der a)—e) für alle Parallelepipede gilt, so gilt es für alle meßbaren $\mathcal{A}$ von endlichem Maß.

Beides wurde a. a. O. Anm.³ gezeigt, wir werden es weiter unten erneut beweisen. α), β) sind ein weiterer Beleg dafür, daß die wichtigsten physikalischen Eigenschaften isomorphieinvariant sind (ohne Rücksicht auf Stetigkeitsfragen!)

¹⁸ D. i. die Zurückführung der kinetischen Gastheorie auf die Maxwell-Boltzmannsche statistische Mechanik, bzw. die Gibbs'sche „Mikrokanonische Gesamtheit“. Vgl. a. a. O. Anm.¹⁷.

¹⁹ Siehe a. a. O. Anm.³. Unmittelbar nachher bewies Herr G. D. Birkhoff mit einer anderen Methode, daß unter denselben Bedingungen sogar a) gilt, Proc. Nat. Ac. 17, Dec. 1931, S.

3. Der Grundgedanke der Koopmanschen Methode ist dieser: Die meßbaren (komplexwertigen) Funktionen $f(P)$ in Ω mit endlichem $\int_{\Omega} |f(P)|^2 dv_P$ bilden einen Hilbertschen Raum $\mathfrak{H}$, falls das innere Produkt (f, g) durch $\int_{\Omega} f(P) \overline{g(P)} dv_P$ definiert wird (also der Betrag $\|f\| = \sqrt{(f, f)}$ durch $\sqrt{\int_{\Omega} |f(P)|^2 dv_P}$)²⁰. Der Funktionaloperator $U_t: f(P) \rightarrow f(S_t P) = U_t f(P)$ ist stetig, und wegen D'_2 (Maßtreue von S_t) unitär²¹. Wegen $S_t S_s \cong S_{t+s}$ ist $U_t U_s = U_{t+s}$. Und aus der (in D'_2 übernommenen) letzten Forderung von D_2 folgt (vgl. § III 2): $(U_t f, g) = \int_{\Omega} f(S_t P) \overline{g(P)} dv_P$ ist eine meßbare Funktion von t . Also:

D_4 . Die U_t sind unitäre Operatoren, die eine einparametrische Gruppe bilden:
 $U_t U_s = U_{t+s}$. U_t hängt meßbar von t ab, d. h. jedes $(U_t f, g)$ ist eine meßbare Funktion von t .²²

Wenn U_t stetig von t abhängt, ist ein Satz von Herrn M. L. Stone anwendbar²³, wonach ein hypermaximaler Operator A mit der Zerlegung der Einheit $E(\lambda)$ existiert²⁴, so daß

$$(U_t f, g) = \int e^{it\lambda} d(E(\lambda) f, g)$$

ist (Stieltjessches Integral!), oder symbolisch:

$$U_t = e^{itA} dE(\lambda) = e^{itA}.$$

Wie der Verfasser a. a. O. Anm. ²² bewies, ist die Stetigkeit nicht erforderlich, es genügt vielmehr hierzu die durch D_4 gewährleistete meßbare Abhängigkeit von t : aus dieser folgt die Stetigkeit und die Gültigkeit des Stoneschen Satzes (vgl. Satz 1, 2, a. a. O.).

Die Aussagen b)–e) können nun alle mit Hilfe der U_t ausgedrückt werden, und ihre weitere Diskussion erfolgt dann auf Grund der Stoneschen Darstellung der U_t : $U_t = e^{itA}$ (vgl. a. a. O. Anm. ^{2, 3} sowie im weiteren Verlauf dieser Abhandlung). Bei den bisherigen Untersuchungen war es dabei so, daß nur die in D_4 angegebenen Eigenschaften der U_t , die der Stoneschen Darstellung gleichwertig sind, eine Rolle spielten, nur bei d)

²⁰ Die Abhandlungen des Verfassers Math. Ann. 102/1, 1929, S. 49, und Math. Ann. 102/3, 1929, S. 370, werden mehrfach zitiert werden, u. zw. als E bzw. A. Zum Obigen vgl. etwa E, S. 108–111 (Anhang I), es handelt sich eigentlich um den Fischer-Rieszschen Satz.

²¹ Es ist $\int_{\Omega} |f(P)|^2 dv_P = \int_{\Omega} |f(S_t P)|^2 dv_P$, d. h. $\|f\| = \|U_t f\|$.

²² Vgl. die zweit-vorhergehende Abhandlung des Verfassers in diesem Bande.

²³ Proc. Nat. Ac. 16, Febr. 1930, S. 172.

²⁴ Vgl. E S. 72, 91, 92.

kommt es noch auf eine feinere Eigenschaft des $E(\lambda)$ -Spektrums an²⁵. Bei tiefergehenden Problemen ist aber zu erwarten, daß die spezielle Entstehungsweise der U_t (d. i. $U_t f(P) = f(S_t P)$) eine Rolle spielen wird. Dieselbe hat nämlich weitere, in D_4 nicht erwähnte Eigenschaften der U_t zur Folge, wie sie a. a. O. Anm. ² von Koopman angegeben wurden. Eine von ihnen ist

$$D_5. \quad U_t f \cdot U_t g = U_t(f \cdot g),$$

falls $f, g, f \cdot g$ alle zum Hilbertschen Raume gehören (der Beweis ist klar). Diese Eigenschaften haben nun die Konsequenz, daß nicht jede Schar $E(\lambda)$ bzw. nicht jedes A auftreten kann, was für die Probleme von unserem Typus von Belang ist. Insbesondere werden die Möglichkeiten für die Struktur des $E(\lambda)$ -Spektrums eingeschränkt.

Ein weiterer Umstand, der Aufmerksamkeit verdient, ist dieser. Sind Ω', Ω'' im Sinne von D_4 mittels der Abbildung Σ isomorph und $S'_t, S''_t = \Sigma S'_t \Sigma^{-1}$ einander entsprechende allgemeine Strömungen, so gilt für ihre Operatorenscharen U'_t, U''_t analog $U''_t = V U'_t V^{-1}$, falls der unitäre Operator V durch $f(P') \rightarrow f(\Sigma^{-1} P'') = V_t f(P'')$ definiert wird. Daraus folgt (da V unsere Hilbertschen Räume isomorph aufeinander abbildet) $E'(\lambda) = V E''(\lambda) V^{-1}, A' = V A'' V^{-1}$. D. h. alle spektralen Eigenschaften stimmen überein, wie es wegen ihrer offenbaren Isomorphieinvarianz auch nicht anders zu erwarten war. Die spektralen Eigenschaften sind aber auch dann invariant, wenn wir bloß die Transformation $U''_t = V U_t V^{-1}$, mit unitärem V haben, ohne daß die Existenz eines Σ , sowie der obige Zusammenhang zwischen ihm und V , gefordert würde. Ja alle Eigenschaften, die mit Hilfe der U_t formuliert werden können, sind so, wir wollen sie daher „unitär-invariant“ nennen. So sind b)–e) unitär-invariant, a) dagegen zunächst nur isomorphieinvariant. Daß die Koopmansche Methode wirklich alle Wahrscheinlichkeitsfragen der klassischen Mechanik erfaßt, wäre belegt, wenn wir wüßten, daß alle isomorphieinvarianten Eigenschaften auch unitär-invariant sind. D. h. wenn wir jedesmal, wenn zwei Ω', Ω'' mit ihren S'_t, S''_t (nach D_1, D_2) gegeben sind, und für ihre U'_t, U''_t $U''_t = V U'_t V^{-1}$ mit unitärem V gilt, eine eindeutige maßtreue Abbildung Σ von Ω' auf Ω'' finden könnten, für die $S''_t = \Sigma S'_t \Sigma^{-1}$ gilt.²⁶

4. Die Einteilung dieser Abhandlung ist diese: In § I diskutieren wir, wie die Zugehörigkeit zu irgendeinem meßbaren A in Ω physikalisch ent-

²⁵ Nämlich auf den sog. Hellingerschen Typus seines Streckenspektrums. Vgl. die zweite in Anm. ³ genannte Abhandlung. Über den Hellingerschen Typus vgl. Hellinger, Dissertation, Göttingen, 1907, sowie Hahn, Monatshefte, 23, 1912, S. 161.

²⁶ D. h. für $W: f(P') \rightarrow f(\Sigma^{-1} P'') = W f(P'')$ soll $U'_t = W U_t W^{-1}$ sein. $W = V$ wird natürlich nicht verlangt.

scheidbar ist. In § II werden die Gültigkeitsbedingungen der Quasi-Ergodenhypothese diskutiert und gezeigt, wie jede allgemeine Strömung als Vereinigung einer Schar solcher Strömungen dargestellt werden kann, deren jede die Q. E. H. erfüllt — oder, wie es abkürzend heißen soll, ergodisch ist.²⁷ In den folgenden Paragraphen betrachten wir demgemäß nur ergodische Strömungen. § III. handelt von den charakteristischen Eigenschaften der Koopmanschen U_t -Scharen, insbesondere wird gezeigt, daß D_λ charakteristisch ist. In § IV. werden für den Fall, daß U_t ein „reines Punktspektrum“ hat, diejenigen Spektra angegeben, die bei Koopmanschen U_t auftreten können.²⁸ Ferner wird gezeigt, daß in diesem Falle das Spektrum allein die U_t bis auf eine Isomorphie (Σ , vgl. 3.) festlegt, d. h. daß hier die am Ende von 3. formulierte Vermutung gilt. § V. ist den Fällen mit Streckenspektrum gewidmet. Hier sind analoge Sätze zu vermuten, bewiesen kann aber vorläufig nur wenig werden. Trotzdem ist gerade der Fall eines reinen Streckenspektrums der eigentlich interessante. Denn einerseits gilt dann und nur dann e . (vgl. das dazu in 3. Gesagte, dieser Satz wurde in der zweiten der in Anm. ³ genannten Abhandlungen bewiesen), d. h. solche Flüsse haben keine über lange Zeiträume gültigen Eigenschaften (vgl. das in § V hierüber zu Sagende); und andererseits ist zu vermuten, daß Strömungen, die nicht ein reines Streckenspektrum haben, seltene Ausnahmen sind. Diese letztere Vermutung wird in § VI durch Diskussion einer ausgedehnten Klasse zweidimensionaler Flüsse belegt, da diese fast alle solche Spektra haben.

Zum Schluß sei noch auf die interessante Analogie zwischen Koopmans Operatoren $U_t = e^{itA}$ und den Operatoren der Quantenmechanik hingewiesen. Die Schrödingersche Wellenfunktion φ (definiert im Zustandsraume des mechanischen Systems und nicht wie unsere f im Phasenraum!) gehorcht bekanntlich in ihrer Abhängigkeit vom Zeitparameter t der Differentialgleichung $\frac{\hbar}{2\pi i} \frac{\partial}{\partial t} \varphi = H\varphi$. Hier ist $\hbar$ das Plancksche Wirkungsquantum und H der Energieoperator. Hieraus folgt sofort $\varphi = e^{it \frac{2\pi}{\hbar} H} \varphi_{(t=0)}$ ²⁹, so daß hier die unitären Operatoren $\hat{U}_t = e^{it \frac{2\pi}{\hbar} H}$ eine fundamentale Rolle spielen. Die Analogie, die durch das Nebeneinanderstellen von A und

²⁷ § I und der Hauptteil von § II wurden bereits in den Proc. Nat. Ac. 18, Jan. 1932, S. 70, publiziert.

²⁸ Bei all diesen Bedingungen wurde der notwendige Charakter von Koopman, a. a. O. Anm. ², bewiesen. Neu ist der schwerere Nachweis ihrer Hinreichendheit.

²⁹ Dies ist bekanntlich die Lösung der Schrödingerschen Differentialgleichung

$$\frac{\partial}{\partial t} \varphi = \frac{2\pi i}{\hbar} H\varphi.$$

$\frac{2\pi}{h}H$ entsteht, ist auffallend³⁰, und es ist nicht schwer, sie zum Nachweis des stetigen Übergehens der Quantenmechanik in die klassische (für $h \rightarrow 0$) auszubauen. Trotzdem scheinen wesentliche mathematische Unterschiede zwischen diesen Operatorscharen zu bestehen. Denn für ein mechanisches System, das in ein endliches Volumen eingesperrt ist, scheint in der Quantenmechanik stets ein reines Punktspektrum vorzuliegen³¹, während im klassisch-mechanischen Problem das reine Streckenspektrum der allgemeine Fall zu sein scheint (vgl. § VI).

I. Meßbare Mengen und physikalische Messungen.

1. Ω sei ein Raum nach D_1 , sein Gesamtmaß endlich: $\mu(\Omega) > 0$, $< +\infty$. Seine Punkte P mögen die Zustände eines mechanischen Systems repräsentieren (Ω sei etwa ein Phasenraum oder ähnliches). Jede Eigenschaft dieses Systems wird dann durch eine geeignete Teilmenge $\mathcal{A}$ von Ω repräsentiert, die aus den Bildpunkten P aller Zustände besteht, die diese Eigenschaft besitzen. Dem l -Maß $\mu(\mathcal{A})$ in Ω schreiben wir die Bedeutung zu, daß wir für die Gültigkeit einer solchen Eigenschaft, d. h. für die Zugehörigkeit von P zu $\mathcal{A}$ die a priori Wahrscheinlichkeit $\frac{\mu(\mathcal{A})}{\mu(\Omega)}$ ansetzen. Und wenn wir wissen, daß P einer gegebenen Menge M angehört, $\mathcal{A} \subseteq M \subseteq \Omega$, die Wahrscheinlichkeit $\frac{\mu(\mathcal{A})}{\mu(M)}$. (Die auftretenden $\mathcal{A}$, M , N , ... sollen alle meßbar sein.)

Eine physikalische Messung entscheidet immer, ob P einem gewissen $\mathcal{A}$ angehört oder nicht. Die Tatsache, daß nur Messungen von beschränkter Genauigkeit möglich sind, daß sich aber die Genauigkeit beliebig steigern läßt, wird wohl am besten so beschrieben:

Z. Für jedes $n = 1, 2, \dots$ ist Ω in eine endliche oder abzählbar-unendliche Anzahl von meßbaren Mengen $N_1^{(n)}, N_2^{(n)}, \dots$ eingeteilt, von denen je zwei einen Durchschnitt vom Maß 0 haben, und deren Vereinigung Ω ist. Dabei sei die $n+1$ -te Einteilung eine Unterteilung der n -ten, d. h. jedes $N_\nu^{(n)}$ sei genau die Vereinigung einiger $N_{\nu'}^{(n+1)}, N_{\nu''}^{(n+1)}, \dots$. Ferner soll $\delta_n = \text{Maximum der Diameter aller } N_\nu^{(n)}$ ($\nu = 1, 2, \dots$)³² für $n \rightarrow \infty$ gegen Null streben.

³⁰ Eine genauere Überlegung zeigt, daß sie vollkommener wird, wenn man die Differentialgleichung der Wellenfunktion durch diejenige des sog. statistischen Operators ersetzt (vgl. z. B.). Dieselbe ist aber ebenso gebaut und das weiter unten zu Sagende gilt auch für sie.

³¹ Dies findet man z. B. bei van der Waerden, Die gruppentheoretische Methode in der Quantenmechanik, Berlin 1932, S. 5, recht klar formuliert.

³² Der Diameter von $\mathcal{A}$ ist das Maximum aller $\overline{x, y}$, x, y in $\mathcal{A}$.

Unter einer Messung n -ter Genauigkeitsstufe verstehen wir eine, welche entscheidet, welchem der $N_1^{(n)}, N_2^{(n)}, \dots P$ angehört³³.

Es sind Messungen aller Genauigkeitsstufen möglich.

Nehmen wir nun an, wir wollten entscheiden, ob P zu einer gegebenen Menge $\mathcal{A}$ gehört, und nehmen zu diesem Zweck eine Messung von der Genauigkeitsstufe n vor. Dann erfahren wir, zu welchem $N_1^{(n)}, N_2^{(n)}, \dots P$ gehört, und zwar ist die a priori Wahrscheinlichkeit dafür, daß $N_\nu^{(n)}$ herauskommt, $\frac{\mu(N_\nu^{(n)})}{\mu(\Omega)}$. Und ist $N_\nu^{(n)}$ herausgekommen, so bedeutet die Zugehörigkeit zu $\mathcal{A}$ die zu $\mathcal{A} \cdot N_\nu^{(n)}, \mathcal{A} \cdot N_\nu^{(n)} \subseteq N_\nu^{(n)} \subseteq \Omega$. Die Wahrscheinlichkeit der Zugehörigkeit zu $\mathcal{A}$ ist also dann $\frac{\mu(\mathcal{A} \cdot N_\nu^{(n)})}{\mu(N_\nu^{(n)})}$. Sei nun

ein $\varepsilon > 0$ gegeben. Ist $\frac{\mu(\mathcal{A} \cdot N_\nu^{(n)})}{\mu(N_\nu^{(n)})} > 1 - \varepsilon$ oder $< \varepsilon$, so können wir auf die Frage „gehört P zu $\mathcal{A}$?“ antworten „ja“ bzw. „nein“, mit einer Wahrscheinlichkeit $> 1 - \varepsilon$, daß die Antwort richtig ist. Ist aber der obige Bruch $\leq 1 - \varepsilon$, $\geq \varepsilon$, so ist keine Antwort mit solcher Sicherheit möglich. Sei die Menge der ν , für die unser Bruch $> 1 - \varepsilon$ oder $< \varepsilon$ ist, $Z(n, \varepsilon)$. Dann ist die a priori Wahrscheinlichkeit dafür, daß wir eine Antwort auf Grund des Meßresultates geben können, deren Wahrscheinlichkeit richtig zu sein $> 1 - \varepsilon$ ist,

$$\frac{\sum_{\nu \text{ in } Z(n, \varepsilon)} \mu(N_\nu^{(n)})}{\mu(\Omega)} = w(n, \varepsilon).$$

Ist nun ein $\delta > 0$ gegeben, und wird für jedes hinreichend große n $w(n, \varepsilon) > 1 - \delta$, so können wir sagen: ist die Messung hinreichend genau, so wird die a priori Wahrscheinlichkeit dafür, daß wir auf Grund der Messung die Frage „gehört P zu $\mathcal{A}$?“ mit einer Wahrscheinlichkeit $> 1 - \varepsilon$ beantworten können werden, ihrerseits $> 1 - \delta$.

Wenn also zu jedem $\varepsilon > 0, \delta > 0$ ein $n_0 = n_0(\varepsilon, \delta)$ existiert, so daß Obiges für alle $n \geq n_0$ gilt, so können wir sagen: die Frage „gehört P zu $\mathcal{A}$ “ kann mit beliebiger Sicherheit durch Messungen entschieden werden. Wir wollen zeigen, daß diese mathematische Forderung in der Tat erfüllt ist.

2. Nehmen wir an, sie wäre es nicht. Dann gäbe es zwei $\varepsilon > 0, \delta > 0$, so daß für unendlich viele n $w(n, \varepsilon) \leq 1 - \delta$ ist, etwa für $n = n_1, n_2, \dots, (n_1 < n_2 < \dots)$. Die Menge $\Omega - \sum_{\nu \text{ in } Z(n_\varrho, \varepsilon)} N_\nu^{(n_\varrho)}$ heiße M_ϱ , nach

³³ Um wohlbekannte Eigenheiten der wirklichen Messung wiederzugeben, könnten wir diese Entscheidungen auch nur mit gewissen Wahrscheinlichkeiten (die gegen 1 bzw. 0 streben) fällen, u. ä. Dies würde aber das Schlußresultat nicht ändern.

Annahme ist $\mu(M_q) \geq \delta \cdot \mu(\Omega)$, dabei sind alle $M_q \subseteq \Omega$, $\mu(\Omega)$ endlich. Nach dem bekannten Satze von Arzelà³⁴ gilt dann für die Menge $\bar{M}$ aller P , die endlich vielen M_q angehören, auch $\mu(\bar{M}) \geq \delta \cdot \mu(\Omega)$. Für ein P von $\bar{M}$ ist für unendlich viele q P in M_q , d. h. P in einem $N_\nu^{(n)q}$, aber in keinem $N_\nu^{(n)q}$, ν aus $Z(n_q, \epsilon)$. D. h. es ist P in $N_\nu^{(n)q}$, ν nicht in $Z(n_q, \epsilon)$. Somit ist für unendlich viele n P in $N_\nu^{(n)}$, ν nicht in $Z(n, \epsilon)$, also

$$\epsilon \leq \frac{\mu(A \cdot N_\nu^{(n)})}{\mu(N_\nu^{(n)})} \leq 1 - \epsilon.$$

Wenn wir dasjenige ν , für welches P zu $N_\nu^{(n)}$ gehört, $\nu(n, P)$ nennen, so ist demnach $\lim_{n \rightarrow \infty} \frac{\mu(A \cdot N_{\nu(n, P)}^{(n)})}{\mu(N_{\nu(n, P)}^{(n)})}$ nicht vorhanden oder $\neq 0, 1$. Dieser

Limes kann also nicht überall, bis auf eine Menge vom Maße 0 existieren und $= 0, 1$ sein.

Aber nach einem Satze von Lebesgue ist er überall in A bis auf eine Menge vom Maße 0 vorhanden und gleich 1³⁵, und wenn wir denselben Satz auf $\Omega - A$ anwenden, so sehen wir, daß er überall in $\Omega - A$ bis auf eine Menge vom Maße 0 vorhanden und gleich 0 ist. Das steht mit dem vorhin Bewiesenen im Widerspruch.

3. Soweit wir beurteilen können, deckt sich unser Standpunkt — Einführung eines Maßes in Ω , und Beurteilung aller Wahrscheinlichkeiten mit seiner Hilfe — mit dem, was bei physikalischen Messungen wirklich geschieht. Akzeptiert man ihn aber, so ist es unmöglich, der in 2 bewiesenen Möglichkeit, die Zugehörigkeit zu meßbaren Mengen zu entscheiden, auszuweichen. Insbesondere ist es wichtig, daß Mengen vom Maße 0 unbeachtet bleiben müssen, Eigenschaften, die nur in solchen gelten, als unbeobachtbar anzusehen sind.

Daß die Stetigkeit dabei gar keine Rolle spielte, liegt daran, daß es von diesem Standpunkte aus sozusagen keine unstetige Funktionen gibt: nach einem Satze von Lusin ist ja für jede meßbare Funktion und jedes

³⁴ Vgl. z. B. de la Vallée Poussin, Cours d'Analyse Infinitésimale, Paris 1909, S. 68—69.

³⁵ Ein Beweis, dem ohne wesentliche Änderungen die hier erforderliche Allgemeinheit zukommt, findet sich bei Carathéodory, „Reelle Funktionen“, S. 492—497, Satz 3. Die dort auftretende Funktion $f(P)$ ist $\begin{cases} = 1, & \text{für } P \text{ in } A \\ = 0, & \text{sonst} \end{cases}$ zu setzen. Der Beweis a. a. O.

beruht auf dem „Überdeckungssatz von Vitali“, S. 299—307. Für die dortigen $\sigma_n(P)$ sind die $N_{\nu(n, P)}^{(n)}$ einzusetzen, da in Ω keine Parallelepipede definiert wurden, brauchen die dortigen Prämissen bez. des Verhältnisses der $\sigma_n(P)$ zu Parallelepipeden nicht zu gelten. Da aber je zwei $N_{\nu(n, P)}^{(n)}$ (n, P variabel!) fremd sind, oder eines das andere umfaßt, passen sie doch ins Gefüge des Beweises. Sogar etwas besser, als die $\sigma_n(P)$, denn die auf S. 301 ff. erfolgende Einführung der Würfel von dreifacher Kantenlänge wird überflüssig.

$\delta > 0$ eine Menge vom Maße $< \delta$ angebbar, außerhalb welcher die Funktion gleichmäßig stetig ist³⁶.

II. Ergodische Strömungen.

1. Ω sei ein Raum nach D_1 , S_t eine allgemeine Strömung in ihm (D_2'), U_t die Koopmansche Operatorenschar, $U_t = e^{itA}$, A hypermaximal, $E(\lambda)$ seine Zerlegung der Einheit (Einleitung 3).

Wegen der Unitarität von U_t und der Schwarzschen Ungleichheit ist

$$|(U_t f, g)| \leq \|U_t f\| \cdot \|g\| = \|f\| \cdot \|g\|,$$

also auch

$$\left| \frac{1}{s-t} \int_t^s (U_\tau f, g) d\tau \right| \leq \|f\| \cdot \|g\|.$$

Infolgedessen kann ein Operator $X_{t,s}$ durch

$$(X_{t,s} f, g) = \frac{1}{s-t} \int_t^s (U_\tau f, g) d\tau \quad (t < s)$$

definiert werden³⁷.

Für jedes meßbare $\mathcal{A}$ endlichen Maßes können wir $\chi_{\mathcal{A}}(P) \begin{cases} = 1, & \text{für } P \text{ in } \mathcal{A} \\ = 0, & \text{sonst} \end{cases}$ bilden, und es gehört zum Hilbertschen Raume. Es wird nun einerseits

$$(X_{t,s} \chi_{\mathcal{A}}, \chi_M) = \int_M X_{t,s} \chi_{\mathcal{A}}(P) dv_P,$$

und andererseits

$$\begin{aligned} (X_{t,s} \chi_{\mathcal{A}}, \chi_M) &= \frac{1}{s-t} \int_t^s (U_\tau \chi_{\mathcal{A}}, \chi_M) d\tau \\ &= \frac{1}{s-t} \int_t^s \left(\int_M \chi_{S_{s-\tau} \mathcal{A}}(P) dv_P \right) d\tau, \end{aligned}$$

also nach dem Satz von Fubini (vgl. Anm. ¹⁰)

$$= \frac{1}{s-t} \int_M \left(\int_t^s \chi_{S_{s-\tau} \mathcal{A}}(P) d\tau \right) dv_P = \frac{1}{s-t} \int_M \Xi(\mathcal{A}, P, t, s) dv_P.$$

Da dies für alle M gilt, ist $X_{t,s} \chi_{\mathcal{A}}(P) = \frac{1}{s-t} \Xi(\mathcal{A}, P, t, s)$ ³⁸.

³⁶ C. R. 154, 1912, S. 1688, oder Sierpinski, Fund. Math. 3, 1922, S. 320. Die dort für das gewöhnliche Lebesguesche Maß geführten Beweise sind für alle l -Maße gültig.

³⁷ Vgl. a. a. O. Anm. ²⁷, S. 72, oder den analogen Beweis in E, S. 112 oben. Die zu verwendenden Bezeichnungen sind auch nach E.

³⁸ Eigentlich braucht dies bloß bis auf eine P -Menge vom Maße 0 zu gelten, aber $X_{t,s} \chi_{\mathcal{A}}$ ist als Element des Hilbertschen Raumes ohnehin nur bis auf eine solche definiert.

Wir berechnen nun $\|X_{t,s}f\|$. Es ist:

$$\begin{aligned}
 \|X_{t,s}f\|^2 &= \|(X_{t,s}f, X_{t,s}f)\| = \frac{1}{s-t} \int_t^s (U_\tau f, X_{t,s}f) d\tau \\
 &= \frac{1}{(s-t)^2} \int_t^s \int_t^s (U_\tau f, U_\sigma f) d\tau d\sigma \\
 &=^{39} \frac{1}{(s-t)^2} \int_t^s \int_t^s (U_{\tau-\sigma} f, f) d\tau d\sigma \\
 &=^{40} \frac{1}{(s-t)^2} \int_{-(s-t)}^{+(s-t)} \int_{2t-|x|}^{2s+|x|} (U_x f, f) \frac{dx dy}{2} \\
 &= \frac{1}{(s-t)^2} \int_{-(s-t)}^{+(s-t)} (s-t-|x|) (U_x f, f) dx.
 \end{aligned}$$

Nun setzen wir nach Einleitung 3 die $E(\lambda)$ für die U_x ein und bemerken, daß die Integrationsreihenfolgen-Vertauschung wegen der absoluten Konvergenz aller auftretenden Integrale zulässig ist. So wird fortlaufend:

$$\begin{aligned}
 &= \frac{1}{(s-t)^2} \int_{-(s-t)}^{+(s-t)} \int_{-\infty}^{+\infty} (s-t-|x|) e^{i\lambda x} dx d(E(\lambda)f, f) \\
 &= \frac{1}{(s-t)^2} \int_{-\infty}^{+\infty} \left(\int_{-(s-t)}^{+(s-t)} (s-t-|x|) e^{i\lambda x} dx \right) d(E(\lambda)f, f) \\
 &= \frac{2}{(s-t)^2} \int_{-\infty}^{+\infty} \left(\int_0^{s-t} (s-t-x) \cos(\lambda x) dx \right) d(E(\lambda)f, f) \\
 &= \frac{2}{(s-t)^2} \int_{-\infty}^{+\infty} \frac{1 - \cos(s-t)\lambda}{\lambda^2} d(E(\lambda)f, f) \\
 &= \int_{-\infty}^{+\infty} \left(\frac{\sin \frac{s-t}{2} \lambda}{\frac{s-t}{2} \lambda} \right)^2 d(E(\lambda)f, f).
 \end{aligned}$$

Der Integrand ist ≥ 0 und der Ausdruck hinter dem d -Zeichen monoton nichtfallend, also können wir das Integral folgendermaßen nach oben abschätzen: Es ist gleich $\int_{-\varepsilon}^{-\varepsilon} + \int_{+\varepsilon}^{+\infty}$ plus $\int_{-\varepsilon}^{+\varepsilon}$. Vergrößern wir den Integranden hierin zu $\frac{1}{\left(\frac{s-t}{2}\varepsilon\right)^2}$ bzw. 1, und ersetzen wir den ersten Integrationsbereich dann durch $\int_{-\infty}^{+\infty}$:

³⁹ Wegen $U_\sigma^* U_\tau = U_\sigma^{-1} U_\tau = U_{\tau-\sigma}$.

⁴⁰ Man substituiere $\sigma - \tau = x$, $\tau + \sigma = y$.

$$\begin{aligned}\|X_{t,s}f\|^2 &\leq \frac{4}{(s-t)^2\epsilon^2} \int_{-\infty}^{+\infty} d(E(\lambda)f, f) + \int_{-\epsilon}^{+\epsilon} d(E(\lambda)f, f) \\ &= \frac{4}{(s-t)^2\epsilon^2} (f, f) + ((E(+\epsilon) - E(-\epsilon))f, f).\end{aligned}$$

Somit ist $\limsup_{s-t \rightarrow +\infty} \|X_{t,s}f\|^2 \leq ((E(+\epsilon) - E(-\epsilon))f, f)$, also wird, wenn für $\epsilon \rightarrow 0$ $(E(+\epsilon) - E(-\epsilon))f \rightarrow 0$, $\lim_{s-t \rightarrow +\infty} \|X_{t,s}f\|^2 = 0$, d. h. $X_{t,s}f \rightarrow 0$ für $s - t \rightarrow +\infty$ ⁴¹.

Als Projektionsoperator, der monoton nichtfallend von ϵ abhängt, hat $E(+\epsilon) - E(-\epsilon)$ einen Limes für $\epsilon \rightarrow 0$: E_0 ⁴². Wir sehen also, daß aus $E_0f = 0$ $X_{t,s}f \rightarrow 0$ für $s - t \rightarrow +\infty$ folgt. Ist hingegen $E_0f = f$, so ist wegen $E(+\epsilon) - E(-\epsilon) \geq E_0$ $E(+\epsilon)f = f$, $E(-\epsilon)f = 0$, d. h. $E(\lambda)f \begin{cases} = f, & \text{für } \lambda > 0 \\ = 0, & \text{für } \lambda < 0 \end{cases}$. Also für alle g $(U_t f, g) = \int_{-\infty}^{+\infty} e^{it\lambda} (E(\lambda)f, g) = (f, g)$, d. h. $U_t f = f$.

Sei $\mathfrak{M}$ die zu E_0 gehörige abgeschlossene Linearmannigfaltigkeit. In ihr gilt stets $E_0f = f$, also $U_t f = f$, also $X_{t,s}f = f$. Für zu $\mathfrak{M}$ orthogonale f gilt $E_0f = 0$, also $X_{t,s}f \rightarrow 0$ ($s - t \rightarrow +\infty$). Somit ist für beide Arten von f $X_{t,s}f \rightarrow E_0f$ ($s - t \rightarrow +\infty$), also auch für alle $f = f_1 + f_2$, f_1 aus $\mathfrak{M}$, f_2 orthogonal zu $\mathfrak{M}$ — d. i. aber jedes f . Wir haben also im Sinne der starken Operatorenkonvergenz $X_{t,s} \rightarrow E_0$ für $s - t \rightarrow +\infty$ ⁴³.

Da aus $E_0f = f$ $U_t f = f$ für alle t folgt, und aus $U_t f = f$ $X_{t,s}f = f$, also $E_0f = f$, können wir $\mathfrak{M}$ auch so charakterisieren: es ist die Menge der gemeinsamen Lösungen aller $U_t f = f$.

2. b), c) aus Einleitung 2 besagen, wie wir nunmehr wissen, dieses: $X_{t,s}\chi_A \rightarrow (\chi_A, \varphi_0) \cdot \varphi_0$ bzw. $(X_{t,s}\chi_A, \chi_M) \rightarrow (\chi_A, \varphi_0) \cdot (\varphi_0, \chi_M)$ ($s - t \rightarrow +\infty$),

wobei $\varphi_0 = \varphi_0(P) = \text{konstant}$ ist, und zwar $= \frac{1}{\sqrt{\mu(\Omega)}}$, wenn $\mu(\Omega)$

endlich ist, und $= 0$, wenn dieses unendlich ist (also $\|\varphi_0\| = 1$ bzw. $\varphi_0 = 0$).

Wir können also auch sagen: $X_{t,s}f \rightarrow P_{[\varphi_0]}f$ bzw. $(X_{t,s}f, g) \rightarrow (P_{[\varphi_0]}f, g)$, falls f, g die Form χ_A haben. Dann muß aber diese Gleichung auch für alle Linearaggregate der χ_A gelten und für deren Häufungspunkte⁴⁴, und

⁴¹ Unter $\rightarrow 0$ verstehen wir im Hilbertschen Raume die starke Konvergenz, d. h. $\|\dots\| \rightarrow 0$. Vgl. A S. 378.

⁴² Vgl. hierfür, wie fürs Folgende, die in E, S. 74—78 entwickelte Theorie der Projektionsoperatoren. Gegenwärtig ist Satz 19 dortselbst gemeint sowie S. 91.

⁴³ Vgl. A, S. 381—384.

⁴⁴ Wegen $|(X_{t,s}f, g)| \leq \|f\| \cdot \|g\|$ sind die $X_{t,s}$ gleichmäßig beschränkt. Vgl. etwa E, S. 73, Satz 12.

dies sind sämtliche f^{45} . Somit besagen b) bzw. c), daß die obigen Gleichungen für alle f, g gelten, d. h., daß $X_{t,s}$ im Sinne der starken bzw. der schwachen Operatorenkonvergenz (vgl. Anm. ⁴³) gegen $P_{[\varphi_0]}$ strebt. Dies gilt übrigens auch dann, wenn b) bzw. c) nicht für alle meßbaren $\mathcal{A}, M$ endlichen Maßes formuliert werden, sondern nur für eine solche Menge von $\mathcal{A}, M$, durch die alle obigen approximiert werden können — etwa alle Summen endlich vieler Kugeln oder Parallelepipede (vgl. a. a. O. Anm. ⁴⁵). Ja, wegen der Additivität dieser Eigenschaften genügen die Parallelepipede selbst.

Andererseits wissen wir aber, daß $X_{t,s} \rightarrow E_0$ (stark). Somit besagen b), c) dasselbe: $E_0 = P_{[\varphi_0]}$, oder, da $\mathfrak{M}$ zu E_0 gehört: $\mathfrak{M} = [\varphi_0]$. Oder: die gemeinsamen Lösungen der $U_t f = f$ sind die $c\varphi_0$, d. h. die Konstanten und nur diese⁴⁶.

Daß $f(P)$ Lösung aller $U_t f = f$ ist, bedeutet: für jedes t gilt $f(S_t P) = f(P)$ bis auf eine P -Menge vom Maße 0. Diese f sind die „Integrale“ der Strömung. Integrale mit endlichem Absolutwertquadrat-Integral, d. h. im Hilbertschen Raume, nennen wir „ H -Integrale“. Wir haben also gezeigt:

SATZ 1. b), c) aus Einleitung 2 sind einander gleichwertig, ferner ist es dasselbe, ob man sie für alle meßbaren $\mathcal{A}, M$ endlichen Maßes formuliert, oder nur für Parallelepipede, oder Summen endlich vieler Kugeln, o. ä. Für ihre Gültigkeit, d. h. für die Ergodizität der Strömung, ist notwendig und hinreichend, daß ihre einzigen H -Integrale die Konstanten seien.

Damit ist der „Ergodensatz“ der klassischen Mechanik bewiesen (vgl. dazu Anm. ³, ¹⁹).

Es ist bei ergodischen Ω noch zweckmäßig zu unterscheiden, ob $\mu(\Omega)$ endlich oder unendlich ist, d. h. ob der Limes von $Z(\mathcal{A}, P, t, s)$ nicht identisch Null ist, oder es ist. Wir nennen Ω eigentlich-ergodisch bzw. Null-ergodisch.

3. Für beliebige, also u. U. nicht ergodische Strömungen können wir die Schlußformel von 1 wieder dazu benutzen, um $\lim_{s-t \rightarrow +\infty} \frac{1}{s-t} Z(\mathcal{A}, P, t, s)$ direkt zu berechnen, dies geschah a. a. O. Anm. ²⁷, S. 75–77. Wir schlagen hier lieber einen indirekten Weg ein, der die in Einleitung 4 angekündigte Zerlegung in ergodische Strömungen mit ergibt.

Ein meßbares $\mathcal{A}$ endlichen Maßes, welches mit jedem $S_t \mathcal{A}$ bis auf eine Menge vom Maße 0 (die natürlich von t abhängt) zusammenfällt, heiße eine $\mathcal{A}$ -Menge. Sie ist auch dadurch charakterisierbar, daß $\chi_{\mathcal{A}} \left(\chi_{\mathcal{A}}(P) \begin{cases} = 1 & \text{für } P \text{ in } \mathcal{A} \\ = 0 & \text{sonst} \end{cases} \right)$ ein H -Integral ist, d. h. zu $\mathfrak{M}$ gehört. Sei

⁴⁵ Vgl. E, S. 109–111.

⁴⁶ Damit die Konstante a zum Hilbertschen Raume gehöre, muß $\int_{\Omega} |a|^2 d\nu_P = |a|^2 \mu(\Omega)$ endlich sein, bei unendlichem $\mu(\Omega)$ also $a = 0$.

umgekehrt f ein beliebiges H -Integral. Es ist starker Häufungspunkt solcher $g(P) = G(f(P))$, deren jede nur endlich vieler verschiedener Werte fähig ist, und diese ihrerseits Linearaggregate endlich vieler solcher $h(P) = H(f(P))$, die nur die Werte 0, 1 annehmen (vgl. a. a. O. Anm. ⁴⁵). Diese sind also $\chi_A(P)$ -Funktionen und ihrer Form nach (H)-Integrale, also sind ihre $\mathcal{A}$ -Mengen. Daher spannen die χ_A aller $\mathcal{A}$ -Mengen die abgeschlossene Linearmannigfaltigkeit $\mathfrak{M}$ auf. Dasselbe gilt, wenn wir uns auf ein Teilsystem der $\mathcal{A}$ -Mengen beschränken, das in diesen überall dicht liegt, und es gibt eine Folge $\mathcal{A}_1, \mathcal{A}_2, \dots$, die dies leistet (vgl. a. a. O. Anm. ⁴⁵). Schließlich können wir diese $\mathcal{A}_n$, durch Änderungen um Mengen vom Maße 0, alle Borelsch machen.

Die Mengenoperationen $\mathcal{A} + M$, $\mathcal{A} \cdot M$, $\mathcal{A} - M$, $\lim_{n \rightarrow \infty} M_n = M_1 \cdot M_2 \cdot \dots$ (mit $M_1 \supset M_2 \supset \dots$) führen aus dem Kreise der Borelschen $\mathcal{A}$ -Mengen nicht hinaus. Infolgedessen kann ein Beweis, den der Verfasser in einem anderen Zusammenhange verwendete, wörtlich übertragen werden⁴⁷, und er ergibt:

Sei A das Intervall $0 \leq x < 1$, $A' = \{x^{(1)}, x^{(2)}, \dots\}$ eine in ihm überall dichte abzählbare Menge. Jedem x von A' ist eine Borelsche $\mathcal{A}$ -Menge $\mathcal{A}(x)$ zugeordnet, derart daß

α) Jedes $\mathcal{A}_n$ entsteht durch die Mengenoperationen $\mathcal{A} + M$, $\mathcal{A} \cdot M$, $\mathcal{A} - M$ aus einigen $\mathcal{A}(x)$.

β) Aus $x \leq y$ (in A') folgt $\mathcal{A}(x) \subseteq \mathcal{A}(y)$ ⁴⁸.

γ) Aus $x_1 > x_2 > \dots > x$, $x_n \rightarrow x$ (für $n \rightarrow \infty$, alles in A') folgt $\lim_{n \rightarrow \infty} \mathcal{A}(x_n) = \mathcal{A}(x_1) \cdot \mathcal{A}(x_2) \cdot \dots = \mathcal{A}(x)$ ⁴⁸.

Sei ferner $K_1, K_2, \dots$ ein „gleichwertiges Umgebungssystem“ (vgl. a. a. O. Anm. ⁴⁵), wir konstruieren ebenso diese Schar:

Jedem ξ von A' ist eine Borelsche Menge $K(\xi)$ zugeordnet, die α), β), γ) ebenfalls erfüllt, jedoch derart, daß durch α) die Ω , $K_1, K_2, \dots$ erzeugt werden.

Wir wollen die Mengensummen und -durchschnitte mit Σ bzw. Π bezeichnen. Es sei

$$\sum_{x \in A'} \mathcal{A}(x) = \Omega_1, \quad \Omega - \Omega_1 = \Omega_2.$$

Da jedes $\mathcal{A}(x) \subseteq \Omega_1$ ist, gilt dasselbe für jedes $\mathcal{A}_n$, also für jede $\mathcal{A}$ -Menge bis auf eine Menge vom Maße 0. Jedes $\mathcal{A}(x)$ ist gegenüber jedem S_t in-

⁴⁷ A, S. 401–403, an Stelle der dortigen Projektionsoperatoren sind $\mathcal{A}$ -Mengen zu setzen, für die $+$, $\cdot$, $-$ ja ebenfalls definiert sind. Insbesondere sind $E_1, E_2, \dots$ durch $\mathcal{A}_1, \mathcal{A}_2, \dots$, die $G(\lambda)$ durch die $\mathcal{A}(x)$, $x = \lambda + 1$, zu ersetzen. A' ist irgendeine in $0 \leq x < 1$ überall dichte Folge, die die dortigen q enthält, etwa die der rationalen Zahlen.

⁴⁸ Man beachte, daß die Konstruktionen von a. O. Anm. ⁴⁷ die Mengenbeziehungen $\subseteq$ und $=$ genauergeben, und nicht bloß bis auf Mengen vom Maße 0.

variant bis auf eine Menge vom Maße 0, also sind es auch Ω_1, Ω_2 . Hätte Ω_2 ein endliches Maß, so wäre es $\mathcal{A}$ -Menge, also da es zu Ω_1 fremd ist, das Maß 0. D. h.: $\mu(\Omega_2) = 0$ oder ∞ . Im ersten Fall brauchen wir Ω_2 nicht zu beachten. Im zweiten ist S_t auch in ihm eine allgemeine Strömung, und jede seiner $\mathcal{A}$ -Mengen hat (da sie fremd zu Ω_1 ist) das Maß 0. Ist f ein H -Integral in Ω_2 , so ist die Menge $|f(P)| > \varepsilon$ ($\varepsilon > 0$) eine $\mathcal{A}$ -Menge, also vom Maße 0, $\varepsilon = 1, \frac{1}{2}, \frac{1}{3}, \dots$ ergibt: $f(P) = 0$ bis auf eine Menge vom Maße 0. Somit ist Ω_2 Null-ergodisch.

Wir definieren nun für jedes P von Ω_1 einen Punkt $[x, \xi]$ im Quadrat $A \times A = \bar{A}$: $0 \leq x, \xi < 1$: $x = \text{finis inf}_{y \text{ in } A'} (P \text{ in } \mathcal{A}(y))$, $\xi = \text{finis inf}_{\eta \text{ in } A'} (P \text{ in } K(\eta))$.

Wenn P, P' $[x, \xi]$ bzw. $[x', \xi']$ entsprechen, so folgt aus $\xi = \xi'$ $P = P'$. Denn durch $\mathcal{A} + M, \mathcal{A} \cdot M, \mathcal{A} - M$ entstehen aus den $K(\eta)$ nur Mengen $(K(\eta_{2\mu}) - K(\eta_{2\mu-1})) + \dots + (K(\eta_{2\nu}) - K(\eta_{2\nu-1}))$ ($\eta_1 < \eta_2 < \dots < \eta_{2\nu-1} < \eta_{2\nu}$, nach β), damit P zu einer solchen gehöre, ist die Zugehörigkeit von ξ zu einem $K(\eta_{2\mu}) - K(\eta_{2\mu-1})$ ($\mu = 1, \dots, \nu$) charakteristisch, d. h. $\eta_{2\nu-1} < \xi < \eta_{2\nu}$. Wegen $\xi = \xi'$ gehören also P, P' zu denselben unter diesen Mengen, also zu denselben K_n , für $P \neq P'$ müßte es aber ein K_n geben, das P enthält und P' nicht. Unsere Abbildung $P \rightarrow [x, \xi]$ ist also eineindeutig, selbst ihre Projektion aus der ξ -Achse ist es noch. Die Abbildung möge Ω_1 auf B abbilden.

Im Produktraume $\Omega_1 \times \bar{A}$ ist die Menge der $[P, x, \xi]$, für die $[x, \xi]$ gerade das Bild von P ist, gleich

$$\prod_{k=2}^{\infty} \sum_{m,n=1}^{k-1} \bar{M}(x_m^{(k)}, x_n^{(k)}; x_{m+1}^{(k)}, x_{n+1}^{(k)}),$$

falls $x_1^{(k)}, \dots, x_k^{(k)}$ die der Größe nach geordneten Zahlen $x^{(1)}, \dots, x^{(k)}$ sind (die k ersten Elemente von A'), und $\bar{M}(y, \eta; z, \zeta)$ die Menge derjenigen $[P, x, \xi]$ ist, für welche entweder $y < x \leq z, \eta < \xi \leq \zeta$ gilt und P in $(\mathcal{A}(z) - \mathcal{A}(y)) \cdot (K(\zeta) - K(\eta))$, oder aber $y < x \leq z, \eta < \xi \leq \zeta$ nicht gilt. Da die $\bar{M}(y, \eta; z, \zeta)$ offenbar Borelsch sind, ist es auch die obengenannte Menge. Ist $M \subseteq \Omega_1$ Borelsch, so ist es auch die Menge der $[P, x, \xi]$ mit P in M , und auch der Durchschnitt jener Menge mit dieser: d. h. die Menge der $[P, x, \xi]$, P in M , $[x, \xi]$ das Bild von P . Ihre Projektion auf den $[x, \xi]$ -Raum $A \times A = \bar{A}$ ist also Suslinsch⁴⁹, dies ist aber das Bild von M . Wir können insbesondere $M = \Omega_1$ setzen, und erhalten: das durch $P \rightarrow [x, \xi]$ vermittelte Bild $\bar{B}$ von Ω_1 , sowie dasjenige jeder Borelschen Teilmenge von Ω_1 , ist eine Suslinsche Menge.

⁴⁹ Vgl. z. B. Hausdorff, „Mengenlehre“, S. 212, Satz 4; oder Lusin, Fund. Math. 10, 1928, S. 1.

Sei $\bar{A}(y, \eta; z, \zeta)$ das Quadrat $y < x \leq z$, $\eta < \xi \leq \zeta$, sein Urbild (d. h. dasjenige seines in $\bar{B}$ liegenden Teiles) ist das meßbare

$$(\mathcal{A}(z) - \mathcal{A}(y)) \cdot (K(\zeta) - K(\eta)),$$

dessen Maß

$$\mu [(\mathcal{A}(z) - \mathcal{A}(y)) \cdot (K(\zeta) - K(\eta))] = \bar{\mu}(y, \eta; z, \zeta).$$

Dies hat die Konsequenz, daß aus

$$\bar{A}(y, \eta; z, \zeta) \cdot \bar{B} \subseteq \sum_1^\infty \bar{A}(y^{(n)}, \eta^{(n)}; z^{(n)}, \zeta^{(n)}) \cdot \bar{B}$$

$$\bar{\mu}(y, \eta; z, \zeta) \leq \sum_1^\infty \bar{\mu}(y^{(n)}, \eta^{(n)}; z^{(n)}, \zeta^{(n)})$$

folgt. Also hat

$$\bar{A}(y, \eta; z, \zeta) \cdot \bar{B} \subseteq \bar{A}(y', \eta'; z', \zeta') \cdot \bar{B}$$

die Folge

$$\bar{\mu}(y, \eta; z, \zeta) \leq \bar{\mu}(y', \eta'; z', \zeta'),$$

und aus $=$ statt $\subseteq$ folgt durch Vertauschung beider Seiten $=$ statt $\leq$. Ferner zieht

$$\bar{A}(y, \eta; z, \zeta) \subseteq \sum_1^\infty \bar{A}(y^{(n)}, \eta^{(n)}; z^{(n)}, \zeta^{(n)})$$

erst recht

$$\bar{\mu}(y, \eta; z, \zeta) \leq \sum_1^\infty \bar{\mu}(y^{(n)}, \eta^{(n)}; z^{(n)}, \zeta^{(n)})$$

nach sich. Schließlich ist $\bar{\mu}(y, \eta; z, \zeta)$ in allen vier Variablen nach rechts halbstetig, weil es $\mathcal{A}(x)$, $K(\xi)$ auch sind.

Wir können nunmehr für die Teilmengen $\bar{C}$ von $\bar{A}$ definieren:

$\bar{\mu}^*(\bar{C})$ sei die untere Grenze aller

$$\sum_1^\infty \bar{\mu}(y^{(n)}, \eta^{(n)}; z^{(n)}, \zeta^{(n)})$$

mit

$$\bar{C} \subseteq \sum_1^\infty \bar{A}(y^{(n)}, \eta^{(n)}; z^{(n)}, \zeta^{(n)}).$$

Aus dem soeben Gezeigten folgt sofort, daß dies ein l -Maß im Quadrat $\bar{A}$ (in bezug auf die gewöhnliche Topologie von $\bar{A}$) ist. Ferner sehen wir noch: $\bar{C}$ und $\bar{C} \cdot \bar{B}$ haben stets dasselbe Maß. Da $\bar{B}$ suslinsch ist, ist es meßbar⁵⁰, also ist bei endlichem $\bar{\mu}^*(\bar{C})$ jedenfalls

⁵⁰ Vgl. Lusin, a. a. O. Anm. ⁴⁹, sein für das gewöhnliche Lebesguesche Maß geführter Beweis gilt für jedes l -Maß.

$$\mu^*(\bar{C} \cdot (\bar{A} - \bar{B})) = \mu^*(\bar{C} - \bar{C} \cdot \bar{B}) = \mu^*(\bar{C}) - \mu^*(\bar{C} \cdot \bar{B}) = 0,$$

d. h.

$$\mu(\bar{C} \cdot (\bar{A} - \bar{B})) = 0.$$

Also hat $\bar{A} - \bar{B}$ selbst das Maß 0.

$[x, \xi] \rightarrow P$ bildet $\bar{B}$ eineindeutig auf Ω_1 ab. Wenn $\bar{C} = \bar{A}(y, \eta; z, \zeta)$ ist, so ist das Urbild von $\bar{C} \cdot \bar{B}$, $(\mathcal{A}(z) - \mathcal{A}(y)) \cdot (K(\zeta) - K(\eta))$ meßbar, und sein μ -Maß ist $\bar{\mu}(y, \eta; z, \zeta)$, d. h. $\bar{\mu}(\bar{C} \cdot \bar{B})$. Dasselbe gilt also, wenn $\bar{C}$ die Summe endlich vieler fremder $\bar{A}(y, \eta; z, \zeta)$ ist, oder Limes einer aufsteigenden Folge solcher, d. h. für jedes offene $\bar{C}$. Von diesen überträgt es sich wiederum durch Limesbildung auf- oder absteigender Folgen auf alle Borelschen $\bar{C}$. Hat $\bar{C}$ das $\bar{\mu}$ -Maß 0, so ist es Teil eines ebensolchen Borelschen $\bar{C}'$; das Urbild von $\bar{C}'$, also auch dasjenige von $\bar{C}$, hat dann auch das μ -Maß 0. Ist schließlich $\bar{C}$ bloß $\bar{\mu}$ -meßbar, so ist es Summe einer Borelschen Menge und einer vom $\bar{\mu}$ -Maß 0; da die obige Behauptung für beide gilt, gilt sie auch für $\bar{C}$.

Für die $\bar{C} \subseteq \bar{B}$ haben wir insbesondere gewonnen: Ist $\bar{C}$ $\bar{\mu}$ -meßbar, so ist sein $[x, \xi] \rightarrow P$ -Bild μ -meßbar und vom selben Maß. Aus Satz 2 der vorhergehenden Abhandlung des Verfassers folgt alsdann, daß $[x, \xi] \rightarrow P$ sowie $P \rightarrow [x, \xi]$ überhaupt maßtreu sind.

4. Das $\bar{\mu}$ -Maß in $\bar{A}$ (oder $\bar{B}$) läßt eine Integraldarstellung zu, die wir angeben wollen.

Wir definieren für alle x, ξ in $0 \leq x \leq 1, 0 \leq \xi \leq 1$ (nicht nur für die in $\bar{A}$!)

$$\bar{\mu}(x, \xi) = \bar{\mu}(\bar{A}(0, 0; x, \xi)).$$

Wenn x, ξ zu A' gehören, ist also

$$\bar{\mu}(x, \xi) = \bar{\mu}(0, 0; x, \xi) = \mu(\mathcal{A}(x) \cdot K(\xi)).$$

Offenbar sind $\bar{\mu}(x, \xi)$ und $\bar{\mu}(x, \xi'') - \bar{\mu}(x, \xi')$ ($\xi' \leq \xi''$) monoton nichtfallend und nach rechts halbstetig in x . Wenn wir also ξ festhalten, x von 0 bis 1 variieren lassen, aber $u = \bar{\mu}(x, 1)$ als unabhängige Variable ansehen, so liegen die Differenzenquotienten von $\bar{\mu}(x, \xi)$ stets im Intervalle 0, 1. Insbesondere ist in den „Konstanzintervallen“ von $u = \bar{\mu}(x, 1)$ $\bar{\mu}(x, \xi)$ ebenfalls konstant. Die von $\bar{\mu}(0, 1) = 0$ bis $\bar{\mu}(1, 1) = \mu(\Omega_1)$ variierende Variable u überspringt allerdings die den Unstetigkeitsstellen x von $\bar{\mu}(x, 1)$ entsprechenden „Sprungintervalle“ $\bar{\mu}(x-0, 1) \leq u < \bar{\mu}(x+0, 1) = \bar{\mu}(x, 1)$, aber in jedem dieser Intervalle können wir $\bar{\mu}(x, \xi)$ linear und an den Enden stetig in u neu definieren. Nimmehr ist der Satz von Lebesgue anwendbar⁵¹: die Derivierte

⁵¹ Vgl. z. B. Carathéodory, „Reelle Funktionen“, S. 551 und S. 545.

$$\varrho(x, \xi) = \lim_{\substack{\varepsilon > 0, \delta > 0 \\ \varepsilon, \delta \rightarrow 0}} \frac{\bar{\mu}(x + \varepsilon, \xi) - \bar{\mu}(x - \delta, \xi)}{\bar{\mu}(x + \varepsilon, 1) - \bar{\mu}(x - \delta, 1)}$$

existiert überall, abgesehen von einer solchen x -Menge, deren u -Menge das gewöhnliche Lebesguesche Maß 0 hat. Die obengenannten „Sprungintervalle“ von u brauchen wir nicht zu beachten, denn in den entsprechenden Unstetigkeitsstellen x existiert der obige Limes, da sein Nenner nicht gegen 0 strebt, und im u -Intervalle hat wegen der Linearität die Derivierte denselben Wert. Wenn wir also für Teilmengen X der x -Menge $0 \leq x \leq 1$ ein Maß

$$\tilde{\mu}^{(0)*}(X) = \int_X d\bar{\mu}(x, 1)$$

einführen⁵², das offenbar ein l -Maß ist, so können wir sagen: $\varrho(x, \xi)$ existiert für alle x , außer für eine Menge vom $\tilde{\mu}^{(0)}$ -Maße 0.

Weiter ist $\bar{w}(x, \xi)$ das Integral seiner Derivierten:

$$\bar{w}(x, \xi) = \int_0^x \varrho(y, \xi) d_y \bar{w}(y, 1).$$

Wir lassen nun ξ nur in A' variieren. Alle diese $\varrho(x, \xi)$ existieren gleichzeitig, bis auf eine x -Menge Y_0 vom Maße 0, deren Komplementärmenge in $0 \leq x \leq 1$ X_0 heiße. Da für $\xi' \leq \xi''$ $\varrho(x, \xi'') - \varrho(x, \xi')$ die Derivierte von $\bar{w}(x, \xi'') - \bar{w}(x, \xi')$ ist, also ≥ 0 , ist $\varrho(x, \xi)$ in ξ monoton nichtfallend. Wir definieren nun für alle ξ

$$\varrho_1(x, \xi) = \lim_{\substack{\eta \text{ in } A', \eta > \xi \\ \eta \rightarrow \xi}} \varrho(x, \eta),$$

$\varrho_1(x, \xi)$ ist in ξ monoton nichtfallend und nach rechts halbstetig. Da $\bar{w}(x, \xi)$ in ξ nach rechts halbstetig ist, wird aus unserer obigen Integralformel durch denselben Grenzprozeß

$$\bar{w}(x, \xi) = \int_0^x \varrho_1(y, \xi) d_y \bar{w}(y, 1).$$

Wir definieren nun jeden Parameterwert x aus X_0 und für jede Teilmenge Ξ der Menge $0 \leq \xi \leq 1$ ein Maß $\tilde{\mu}_x(\Xi) = \int_{\Xi} d\xi \varrho_1(x, \xi)$ (vgl. das am Ende von Anm. ⁵² Gesagte), das offenbar ein l -Maß ist. Dabei ist die gewöhnliche Topologie von $0 \leq \xi \leq 1$ zugrunde zu legen. Wenn $\bar{C}$ eine Menge $\subseteq \bar{A}$ ist, so sei $(\bar{C})_{(x)}$ die Menge aller ξ mit $[x, \xi]$ in $\bar{C}$. Sei $\bar{C}$ die Menge $\bar{A}(0, 0; x, \xi)$, dann ist:

⁵² D. i. das äußere Maß der entsprechenden u -Menge, wenn diese unmeßbar ist, stellt $\int_X$ das „obere Lebesguesche Integral“ dar.

$$\begin{aligned}
\bar{\mu}(\bar{C}) &= \bar{w}(x, \xi) = \int_0^x \varrho_1(y, \xi) dy \bar{w}(y, 1) \\
&= \int_0^x \left[\int_0^\xi d\eta \varrho_1(y, \eta) \right] dy \bar{w}(y, 1) = \int_0^1 \left[\int_{(\bar{C})_{(y)}} d\eta \varrho_1(y, \eta) \right] dy \bar{w}(y, 1) \\
&= \int_0^1 \tilde{\mu}_{(y)}((\bar{C})_{(y)}) dy \bar{w}(y, 1).
\end{aligned}$$

Wir fragen allgemein: für welche $\bar{C} \subseteq \bar{A}$ gilt die Formel

$$\bar{\mu}(\bar{C}) = \int_0^1 \tilde{\mu}_{(y)}((\bar{C})_{(y)}) dy \bar{w}(y, 1),$$

wobei auch verlangt wird, daß der Integrand überall, bis auf eine y -Menge vom $\tilde{\mu}^{(0)}$ -Maße 0, Sinn hat (nicht bloß als $\tilde{\mu}_{(y)}^*$, sondern als $\tilde{\mu}_{(y)}$, d. h. $(\bar{C})_{(y)}$ soll $\tilde{\mu}_{(y)}$ -meßbar sein), und eine $\tilde{\mu}^{(0)}$ -meßbare Funktion ist (d. h. das Integral Sinn hat)? Dies ist für die $\bar{C} = \bar{A}(0, 0; x, \xi)$ der Fall, also auch für ihre Differenzen, die $\bar{C} = \bar{A}(x, \xi; y, \eta)$; sodann für die Summen endlich vieler solcher und die Limites aufsteigender Folgen solcher Mengen — d. i. für alle offenen $\bar{C}$. Genau wie in der Betrachtung am Ende von 3. überträgt sich unsere Formel von diesen $\bar{C}$ successiv auf die Borelschen $\bar{C}$, auf die $\bar{C}$ vom $\bar{\mu}$ -Maße 0, und schließlich auf alle $\bar{\mu}$ -meßbaren $\bar{C}$.

Zum Schluß stellen wir noch dieses fest. Sei X eine Menge in A , $\bar{X}$ die Menge aller $[x, \xi]$, x in X , $0 \leq \xi < 1$ (in $\bar{A}$). Wenn $\bar{X}$ $\bar{\mu}$ -meßbar ist, ergibt unsere Formel sofort: X ist $\tilde{\mu}^{(0)}$ -meßbar, $\tilde{\mu}^{(0)}(X) = \bar{\mu}(\bar{X})$. Ist umgekehrt X $\tilde{\mu}^{(0)}$ -meßbar, so gibt es zwei Borelsche X', X'' mit $X' \subseteq X \subseteq X''$, $\tilde{\mu}^{(0)}(X'' - X') = 0$, also sind $\bar{X}', \bar{X}''$ auch Borelsch, also meßbar, und es ist $\bar{\mu}(\bar{X}'' - \bar{X}') = 0$. Wegen $\bar{X}' \subseteq \bar{X} \subseteq \bar{X}''$ ist auch $\bar{X}$ meßbar und hat dasselbe $\bar{\mu}$ -Maß wie $\bar{X}, \bar{X}''$, d. h. wie X', X'' , d. h. wie X . Also: $X, \bar{X}$ sind immer gleichzeitig meßbar und haben dasselbe ($\tilde{\mu}^{(0)}$ - bzw. $\bar{\mu}$ -)Maß. Da $\bar{A} - \bar{B}$ das Maß 0 hat, gilt dasselbe, wenn wir $\bar{X}$ auf $\bar{B}$ einschränken.

5. Wir übertragen nun $\bar{B}$ durch die maßtreue Abbildung $[x, \xi] \rightarrow P$ zurück nach Ω_1 . Aus der Menge der $[y, \eta]$ mit $y \leq x$, $0 \leq \eta \leq 1$ wird $A(x)$, aus der mit $y = x$, $0 \leq \eta \leq 1$ also

$$M(x) = \prod_{\substack{y \text{ in } A' \\ y > x}} A(y) - \sum_{\substack{y \text{ in } A' \\ y < x}} A(y).$$

Wir übertragen dabei auch die in $\bar{A}$ und seinen Teilmengen definierten l -Maße: aus $\bar{\mu}$ in $\bar{B}$ wird, wie wir wissen, μ in Ω_1 ; $\tilde{\mu}_{(x)}$ übertragen wir auf $M(x)$, wo es $\mu_{(x)}$ heißen soll — all dies sind l -Maße. Immerhin ist bei $\mu_{(x)}$ dieses zu beachten: Bei ihm ist jene Topologie zugrunde zu

legen, die bei $\xi \rightarrow$ ins zu $[x, \xi]$ gehörige P aus der gewöhnlichen Topologie von $0 \leq \xi \leq 1$ entsteht. Eine andere denkbare Topologie für $M(x)$ ist die, bei der die in $M(x)$ relativ-offenen Mengen (im Sinne der Ω -Topologie, also die $O \cdot M(x)$, O offen in Ω) als offen gelten. Nun sind die Borelschen Mengen der ersteren Topologie aus den Bildern der ξ -Mengen $\xi \leq \eta$, η in A' erzeugbar, d. h. aus den Mengen $M(x) \cdot K(\eta)$, oder, was nach Konstruktion der $K(\eta)$, aus den $M(x) \cdot K_n$. Die Borelschen Mengen der letzteren Topologie entstehen aus den $M(x) \cdot O$, also auch aus den $M(x) \cdot K_n$. Beide Topologien haben also dieselben Borelschen Mengen, da aber für den l -Maßcharakter nur die Borelschen Mengen maßgebend sind (vgl. Anm. ⁹ in der vorhergehenden Abhandlung des Verfassers), ist $\mu_{(x)}$ auch im Sinne der letzteren Topologie, d. i. der auf $M(x)$ übertragenen Topologie von Ω , ein l -Maß.

Ist $M[X] = \sum_{x \in X} M(x)$, so ist $M[X]$ das Bild von $\bar{X}$, also sind auch X und $M[X]$ gleichzeitig meßbar, und haben dieselben ($\tilde{\mu}^{(0)}$ - bzw. μ -)Maße. Aus unserer $\tilde{\mu}(C)$ -Formel aber wird: für jedes meßbare $N \subseteq \Omega_1$ ist

$$\mu(N) = \int_0^1 \mu_{(x)}(M(x) \cdot N) dw(x).$$

Dabei ist $w(x) = w(x, 1) = \tilde{\mu}^{(0)}$ (Menge aller y mit $0 < y \leq x$). Der Integrand kann für einige x sinnlos sein ($\mu_{(x)}^*$ undefiniert oder $M(x) \cdot N$ unmeßbar), aber die Menge dieser x hat das $\tilde{\mu}^{(0)}$ -Maß 0; insbesondere ist $\mu_{(x)}^*$ nur für die x von X_0 definiert, aber diese Komplementärmenge von X_0 in $0 \leq x \leq 1$, Y_0 , hat das $\tilde{\mu}^{(0)}$ -Maß 0.

Untersuchen wir nunmehr das Verhältnis der S_t zu den $M(x)$. Sei T der m -Raum der t , $\tilde{\mu}$ das gewöhnliche Lebesguesche Maß in ihm, $\Omega \times T$ der Produktraum, $\bar{\mu}$ das dort nach D_1 definierte l -Maß (die Verwechslung mit dem $\bar{\mu}$ in 4. droht nicht, da dieses nicht mehr vorkommen wird).

Sei $M \subseteq \Omega$. Wenn es offen ist, ist sein $[P, t] \rightarrow S_t P$ -Urbild in $\Omega \times T$ meßbar (d. h. die Funktion $S_t P$ ist meßbar, vgl. Anm. ¹¹), und von den offenen M überträgt sich dies mühelos auf die Borelschen. Hat M außerdem das Maß 0, so ist für jedes feste t die Menge der P mit $S_t P$ in M vom Maße 0 (S_t ist maßtreu), also ergibt der Satz von Fubini (a. a. O. Anm. ¹⁰), den wir auf das Urbild wegen seiner Meßbarkeit anwenden dürfen: das Urbild selbst hat das Maß 0. Hat M bloß das Maß 0, so ist $M \subseteq M'$, M' Borelsch und vom Maß 0, also Urbild von $M \subseteq$ Urbild von M' , das vom Maße 0 ist, also das Urbild von M auch. Das Urbild von M ist also meßbar, wenn M Borelsch ist, oder vom Maße 0, also auch für die Summen zweier solcher M , d. h. für jedes meßbare M .

Gehöre x zu A' . $A(x)$ ist meßbar, also auch die Menge der $[P, t]$, P in $A(x)$, ferner das $S_t P$ -Urbild: die Menge der $[P, t]$, $S_t P$ in $A(x)$.

Also ist auch die Summe minus der Durchschnitt dieser zwei Mengen meßbar, d. i. die Menge der $[P, t]$, für die P zu $\mathcal{A}(x)$ gehört, aber $S_t P$ nicht, oder umgekehrt. Und auch die Summe dieser Mengen für alle x von A' ist meßbar, sie heiße $\bar{Y}_1$. Nun sei t fest. Die Menge der P , für die $[P, t]$ zu $\bar{Y}_1$ gehört, ist

$$\sum_{x \text{ in } A'} [(\mathcal{A}(x) + S_t^{-1} \mathcal{A}(x)) - (\mathcal{A}(x) \cdot S_t^{-1} \mathcal{A}(x))],$$

also, wie jeder ihrer Addenden, vom Maße 0. Der Satz von Fubini ist anwendbar: $\bar{Y}_1$ hat das Maß 0. $[P, t]$ nicht in $\bar{Y}_1$ bedeutet: P gehört zu denselben $\mathcal{A}(x)$ (x in A'), wie $S_t P$, also auch zu denselben $M(x)$ ($0 \leq x < 1$), oder: P gehört dann und nur dann zu $M(x)$, wenn $S_t P$ es tut, u. zw. für alle $0 \leq x < 1$.

Nun definieren wir $S_t P$ um, indem wir es in $\bar{Y}_1$ für undefiniert erklären, sonst aber nicht ändern. Da $\bar{Y}_1$ das Maß 0 hat, bleibt $S_t P$ meßbar; da bei festem t die Menge der P mit $[P, t]$ in $\bar{Y}_1$ das Maß 0 hat, sind die neuen S_t den alten äquivalent. Dabei sind alle $M(x)$ den neuen S_t gegenüber genau invariant. Auf Grund der Ersetzungsmöglichkeit durch eine äquivalente allgemeine Strömung dürfen wir also annehmen, daß die $M(x)$ von vornherein genau invariant waren.

Da jedes $M(x)$ genau S_t -invariant ist, ist es auch jedes $M[X]$. Ist X meßbar, $\tilde{\mu}^{(0)}(X)$ endlich, so ist auch $M[X]$ meßbar, $\mu(M[X])$ endlich, d. h. $M[X]$ $\mathcal{A}$ -Menge. Umgekehrt ist jedes $\mathcal{A}(\eta)$ ein $M[X]$ (X ist die Menge $0 \leq \xi \leq \eta$) mit meßbarem X ; also auch alles, was durch $\mathcal{A} + M$, $\mathcal{A} \cdot M$, $\mathcal{A} - M$ daraus entsteht, d. h. jedes $\mathcal{A}_n$; also auch alles, was durch Limesbildung an auf- oder absteigenden Folgen daraus entsteht, d. h. jede $\mathcal{A}$ -Menge (bis auf eine Menge vom Maße 0, vgl. 3. und a. O. Anm. ⁴⁵). Dabei ist das Maß der $\mathcal{A}$ -Menge, d. h. $\tilde{\mu}^{(0)}(X) = \mu(M[X])$ endlich. Also: alle $M[X]$ mit meßbarem X und endlichem $\tilde{\mu}^{(0)}(X) = \mu(M[X])$ sind, bis auf Mengen vom Maße 0, mit allen $\mathcal{A}$ -Mengen identisch.

6. Betrachten wir die Abbildung $[P, t] \rightarrow [S_t P, t]$. Ihr Definitionsbereich ist derjenige von $S_t P$, also meßbar; für jedes feste t bilden die P , deren $[P, t]$ nicht in ihm liegt, eine Menge vom Maße 0. Seine Komplementärmenge hat also, nach dem Satz von Fubini, auch als $[P, t]$ -Menge das Maß 0. Betrachten wir das Urbild einer Menge $\bar{N}$ (d. i. die Menge der $[P, t]$ mit $[S_t P, t]$ in $\bar{N}$, $\bar{N}$ braucht keineswegs aus lauter $[S_t P, t]$ -s zu bestehen!). Ist $\bar{N} = O \times \tilde{O}$, $O, \tilde{O}$ Borelsch, so ist dieses Urbild der Durchschnitt des $[P, t] \rightarrow [S_t P]$ -Urbildes von O , das wegen der Meßbarkeit von $S_t P$ meßbar ist, mit der offenbar meßbaren Menge der $[P, t]$ mit t in $\tilde{O}$ — also meßbar. Bei gegebenem t ist die Menge der P mit $[P, t]$ im Urbild leer, wenn t nicht zu $\tilde{O}$ gehört, und $S_t^{-1} O$, wenn t zu $\tilde{O}$ gehört. Nach dem Satz von Fubini ist also das Maß des Urbildes

$$\begin{aligned}\int_{\tilde{O}} \mu(S_t^{-1} O) dt &= \int_{\tilde{O}} \mu(O) dt = \mu(O) \cdot \tilde{\mu}(\tilde{O}) \\ &= \bar{\mu}(O \times \tilde{O}) = \bar{\mu}(\tilde{N}) \quad (\text{vgl. } D_1).\end{aligned}$$

Die Aussage: das Urbild von $\bar{N}$ ist meßbar und hat dasselbe Maß wie $\bar{N}$, gilt also für $\tilde{N} = O \times \tilde{O}$, $O, \tilde{O}$ Borelsch. Somit gilt sie auch für alle endlichen $O \times \tilde{O}$ -Summen (da die Addenden fremd wählbar sind), für Limes aufsteigender Folgen derselben, d. i. für alle offenen Mengen. Und weiter überträgt sie sich durch Limesbildung an auf- oder absteigenden Folgen auf alle Borelschen $\tilde{N}$. Da jedes $\tilde{N}$ vom Maße 0 Teilmenge eines Borelschen $\tilde{N}'$ vom Maße 0 ist, gilt sie auch für diese, also auch für die Summen Borelscher Mengen und solcher vom Maße 0 -- d. h. für alle meßbaren $\tilde{N}$.

$[P, t] \rightarrow [S_t, P]$ bildet also $\Omega \times T$ minus eine Menge vom Maße 0 eindeutig auf eine gewisse Teilmenge $\tilde{N}^{(0)}$ von $\Omega \times T$ ab, und das Urbild jedes meßbaren $\tilde{N}$ ist meßbar und vom selben Maße wie $\tilde{N}$ (das keineswegs $\subseteq \tilde{N}^{(0)}$ zu sein braucht). Betrachtet man den Beweis von Satz 2 in der vorhergehenden Abhandlung des Verfassers; so erkennt man, daß er auf diesen Fall anwendbar ist, und $\tilde{N}^{(0)} = \Omega \times T$ minus eine Menge vom Maße 0 ergibt. Ist dies gesichert, so kann man den genannten Satz 2 selbst anwenden, aus ihm folgt dann die Maßtreue von $[P, t] \rightarrow [S_t P, t]$. Die inverse Abbildung $[P, t] \rightarrow [S_t^{-1} P, t]$ ist also auch maßtreu.

Sei N meßbar. Ist $\tilde{N}$ die Menge aller $[P, t]$, P in N , so ist es auch maßtreu. Somit ist sein $[S_t P, t]$ - und sein $[S_t^{-1} P, t]$ -Urbild meßbar, d. h. das $[P, t] \rightarrow S_t P$ - und das $[P, t] \rightarrow S_t^{-1} P$ -Urbild von N (aus der bloßen Meßbarkeit von $S_t P$ würde nur das erstere, und nur für offene oder höchstens für Borelsche N folgen). Also ist $S_t^{-1} P$ meßbar als $[P, t]$ -Funktion.

Wir können also $S_t P$ undefinieren: für $t > 0$ sei es das alte, für $t = 0$ die Identität, für $t < 0$ $S_{-t}^{-1} P$. Das neue $S_t P$ ist also meßbar (in $[P, t]$), und nach D_2' dem alten äquivalent, es ist aber jetzt genau (nicht bloß bis auf Mengen vom Maße 0) $S_t S_{-t} = S_{-t} S_t = S_0 = \text{Identität}$. Wir dürfen also annehmen, daß dies von vornherein so war (die Invarianz der $M(x)$ wird dadurch nicht berührt).

Betrachten wir ferner das $S_t S_s P$ -Urbild eines meßbaren N . Zunächst bei festem t : dies ist das $S_s P$ -Urbild von $S_t^{-1} N$, welches meßbar ist, also meßbar. Also ist $S_t S_s P$ als $[P, s]$ -Funktion meßbar. Zweitens sei alles variabel. Dann bilden wir zuerst das $[P, t] \rightarrow S_t P$ -Urbild von $N, \bar{N}'$, welches meßbar ist, und sodann das $[P, t, s] \rightarrow [S_s P, t]$ -Urbild von $\bar{N}', \bar{N}''$, das auch meßbar ist. Letzteres kann genau so bewiesen werden, wie weiter oben das Entsprechende für das $[P, t] \rightarrow [S_t P]$ -Urbild, nur ist an Stelle der maßtreuen Abbildung $P \rightarrow S_t P$ (t fest) die ebenfalls maßtreue

Abbildung $[P, t] \rightarrow [S_s P, t]$ (s fest) als Ausgangspunkt zu nehmen. Das meßbare $\bar{N}''$ ist das $[P, t, s] \rightarrow S_t S_s P$ -Urbild von N . Also ist $S_t S_s P$ auch als $[P, t, s]$ -Funktion meßbar.

Schließlich wollen wir die Formel $\mu(N) = \int_0^1 \mu_{(x)}(M(x) \cdot N) dw(x)$ ($N \subseteq \Omega_1$) auf $\Omega_1 \times T$ und $\Omega_1 \times T \times T$ übertragen. Wenn ein $\bar{N} \subseteq \Omega_1 \times T$ oder ein $\bar{\bar{N}} \subseteq \Omega_1 \times T \times T$ gegeben ist, verstehen wir unter $(\bar{N})_{(t)}$ bzw. $(\bar{\bar{N}})_{(t,s)}$ die Menge aller P mit $[P, t]$ in $\bar{N}$ bzw. $[P, t, s]$ in $\bar{\bar{N}}$ (also $\subseteq \Omega_1$). Die neuen Formeln lauten:

$$\bar{\mu}(\bar{N}) = \int_0^1 \int_{-\infty}^{\infty} \mu_{(x)}(M(x) \cdot (\bar{N})_{(t)}) dw(x) dt,$$

$$\bar{\bar{u}}(\bar{\bar{N}}) = \int_0^1 \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \mu_{(x)}(M(x) \cdot (\bar{\bar{N}})_{(t,s)}) dw(x) dt ds$$

($\bar{N} \subseteq \Omega_1 \times T$, $\bar{\bar{N}} \subseteq \Omega_1 \times T \times T$, der Integrand soll bis auf eine $[x, t]$ - bzw. $[x, t, s]$ -Menge vom Maße 0 sinnvoll sein, und meßbare Funktion von $[x, t]$ bzw. $[x, t, s]$). Für $\bar{N} = O \times \tilde{O}$ bzw. $\bar{\bar{N}} = O \times \tilde{O} \times \tilde{O}'$ (vgl. D_1) folgt dies unmittelbar aus der ursprünglichen Formel, und von diesen überträgt man es, wie beim Maßtreuebeweise w. o., sukzessiv auf alle offenen Borelschen vom Maße 0 und alle meßbaren Mengen $\bar{N}$ bzw. $\bar{\bar{N}}$.

7. Die Menge $\bar{\bar{N}}_2$ der $[P, t, s]$ mit $S_t S_s P \neq S_{t+s} P$ ist meßbar, weil auf beiden Seiten meßbare Funktionen stehen. Bei festem t, s ist die P -Menge mit zu ihr gehörigem $[P, t, s]$ vom Maße 0 (D'_2), also ist sie es selbst auch. Somit ist

$$\int_0^1 \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \mu_{(x)}(M(x) \cdot (\bar{\bar{N}}_2)_{(t,s)}) dw(x) dt ds = 0,$$

$$\text{d. h. } \mu_{(x)}(M(x) \cdot (\bar{\bar{N}}_2)_{(t,s)}) = 0,$$

abgesehen von einer $[x, t, s]$ -Menge, $\bar{\bar{Y}}_2$, vom Maße 0. Die $\mu_{(x)}$ -Bedingung besagt: $S_t S_s P = S_{s+t} P$ für alle P von $M(x)$, bis auf eine Menge vom ($\mu_{(x)}$ -) Maß 0, oder: $S_t S_s \cong S_{t+s}$ in $M(x)$ (D'_2). Hätten wir t von vornherein festgehalten, und wären wir von der Meßbarkeit der $[P, s]$ -Menge mit $S_t S_s P \neq S_{t+s} P$ ausgegangen, so hätten wir ebenso erhalten: abgesehen von einer $[x, s]$ -Menge vom Maße 0, gilt $S_t S_s \cong S_{t+s}$ in $M(x)$, d. h. $[x, t, s]$ gehört zu $\bar{\bar{Y}}_2$.

Wenden wir nun den Satz von Fubini auf $\bar{\bar{Y}}_2$ an: Abgesehen von einer $[x, t]$ -Menge vom Maße 0, $\bar{\bar{Y}}_2$, gehören alle $[x, t, s]$, bis auf eine s -Menge vom Maße 0, zu $\bar{\bar{Y}}_2$. Dasselbe gilt wegen des zuletzt Gesagten auch bei festem t , für alle x bis auf eine Menge vom Maße 0. D. h.: $\bar{\bar{Y}}_2$ hat das Maß 0, und bei festem t hat die x -Menge mit $[x, t]$ in $\bar{\bar{Y}}_2$ auch das Maß 0.

Sei N meßbar, $\bar{N}$ sein $[P, t] \rightarrow S_t P$ -Urbild, also auch meßbar.

$$\mu_{(x)}(M(x) \cdot (\bar{N})_{(t)}) = \mu_{(x)}(M(x) \cdot S_t^{-1} N)$$

ist dann eine meßbare $[x, t]$ -Funktion, ebenso $\mu_{(x)}(M(x) \cdot N)$. Die $[x, t]$ -Menge $\bar{Y}_3(N)$ der $\mu_{(x)}(M(x) \cdot S_t^{-1} N) \neq \mu_{(x)}(M(x) \cdot N)$ ist also meßbar. Wir wenden nun unsere μ -, $\mu_{(x)}$ -Formel auf $M[X] \cdot S_t^{-1} N$ und $M[X] \cdot N$ an, es ist:

$$\begin{aligned} \mu(M[X] \cdot S_t^{-1} N) &= \int_0^1 \mu_{(x)}(M(x) \cdot M[X] \cdot S_t^{-1} N) \, dw(x) \\ &= \int_X \mu_{(x)}(M(x) \cdot S_t^{-1} N) \, dw(x), \\ \mu(M[X] \cdot N) &= \int_0^1 \mu_{(x)}(M(x) \cdot M[X] \cdot N) \, dw(x) \\ &= \int_X \mu_{(x)}(M(x) \cdot N) \, dw(x). \end{aligned}$$

Wegen der Maßtreue von S_t und der Invarianz der $M(x)$, $M[X]$ sind die linken Seiten gleich, also auch die rechten. Dies gilt für alle X , daher ist bis auf eine x -Menge vom Maße 0

$$\mu_{(x)}(M(x) \cdot S_t^{-1} N) = \mu_{(x)}(M(x) \cdot N).$$

D. h.: für jedes feste t ist die x -Menge mit $[x, t]$ in $\bar{Y}_3(N)$ vom Maße 0. Nach dem Satz von Fubini ist daher auch $\bar{Y}_3(N)$ vom Maße 0.

Sei nun $\bar{Y}_3 = \sum_{\xi \text{ in } A'} \bar{Y}_3(K(\xi))$. Auch für $\bar{Y}_3$ gilt: es hat das Maß 0, und bei festem t hat die x -Menge mit $[x, t]$ in $\bar{Y}_3$ auch das Maß 0. Liegt $[y, t]$ nicht in $\bar{Y}_3$, so gilt

$$\mu_{(y)}(S_t^{-1} \cdot M) = \mu_{(y)}(M) \quad (M, S_t^{-1} M \text{ meßbar!})$$

für alle $M = M(y) \cdot K(\eta)$, d. h. für diejenigen Urbilder $M \subseteq M(y)$, die nach der Abbildung $P \rightarrow [x, \xi]$ aus 3 den Strecken $[x, \xi]$ mit $x = y$, $0 \leq \xi \leq \eta$ entsprechen. Hieraus folgt, daß es auch für die Urbilder von Differenzen solcher Strecken gilt, sowie für die von Summen endlich vieler solcher Differenzen und für die Limites aufsteigender Folgen solcher Summen. D. i. für die Urbilder von offenen Mengen, von welchen man auf die gewohnte Weise zu den Urbildern Borelscher, vom Maße 0 ($\hat{\mu}_{(y)}$ -Maß!), diese Mengen haben wegen der maßtreuen Übertragung zu $\mu_{(y)}$, selbst das $\mu_{(x)}$ -Maß 0), und meßbarer Mengen gelangt. Diese letzteren Urbilder sind aber einfach alle ($\mu_{(y)}$ -) meßbaren $M \subseteq M(y)$, unsere Formel gilt mithin auch für diese. Für $M = M(y)$ insbesondere wird demnach

$$\mu_{(y)}(S_t^{-1} M(y)) = \mu_{(y)}(M(y)),$$

und da $S_t^{-1} M(y) \subseteq M(y)$ und $\mu_{(y)}(M(y))$ endlich ist (es ist ≤ 1 , denn $M(y)$ ist der Limes der Urbilder der $\xi < \zeta \leq \eta$, $\xi \rightarrow 0$, $\zeta \rightarrow 1$ (vgl. 3), und diese haben die Maße $\varrho(y, \eta) - \varrho(y, \xi) \leq 1$), ist $M(y) - S_t^{-1} M(y)$ vom Maße 0. D. h. $S_{-t} = S_t^{-1}$ ist überall in $M(y)$ definiert bis auf eine Menge vom $(\mu_{(x)})$ -Maße 0.

Wenn also $[x, t]$, $[x, -t]$ beide nicht in $\bar{Y}_3$ liegen, so bildet $S_t M(x)$ minus eine Menge vom $(\mu_{(x)})$ -Maße 0 auf $M(x)$ minus eine Menge vom $(\mu_{(x)})$ -Maße 0 ab und ist $(\mu_{(x)})$ -maßtreu.

8. Die Menge der $[x, t]$, die in $\bar{Y}_2$ oder $\bar{Y}_3$ liegen oder für die $[x, -t]$ in $\bar{Y}_3$ liegt, heiße $\bar{Y}_4$. $\bar{Y}_4$ hat wieder das Maß 0, und bei festem t hat die Menge der x mit $[x, t]$ in $\bar{Y}_4$, d. i. $(\bar{Y}_4)_{(t)}$, das Maß 0. Gehört $[x, t]$ nicht zu $\bar{Y}_4$, so bildet $S_t M(x)$ bis auf Mengen vom Maße 0 (vgl. am Ende von 7) eineindeutig und maßtreu auf sich selbst ab, und für alle s , bis auf eine Menge vom Maße 0, gilt $S_t S_s \cong S_{t+s}$ in $M(x)$.

Wenn $\bar{X} \subseteq A \times T$ ist, so sei $\bar{M}[\bar{X}]$ die Menge aller $[P, t]$, für die P zu einem $M(x)$ mit $[x, t]$ in $\bar{X}$ gehört. So wie wir am Ende von 4 zeigten: mit X ist auch $M[X]$ meßbar und vom selben Maß ($\bar{\mu}^{(0)}$ bzw. μ), gilt auch: mit $\bar{X}$ ist auch $\bar{M}[\bar{X}]$ meßbar und vom selben Maß ($\bar{\mu}^{(0)}$ bzw. $\bar{\mu}$). Denn für $\bar{X} = X \times \tilde{O}$ (vgl. D_1) folgt dies, aus der Behauptung über die $X, M[X]$, und von diesen überträgt es sich auf die gewohnte Weise sukzessiv auf die offenen, Borelschen, vom Maße 0, und die meßbaren $\bar{X}$.

Also hat $\bar{M}[\bar{Y}_4]$ das $(\bar{\mu})$ -Maß 0, und bei festem t hat $\bar{M}[(\bar{Y}_4)_{(t)}]$, welches offenbar mit der Menge der P mit $[P, t]$ in $\bar{M}[Y_4]$ zusammenfällt, ebenfalls das (μ) -Maß 0. Wenn wir also die Definition von $S_t P$ innerhalb von $\bar{M}[\bar{Y}_4]$ beliebig ändern, wird das neue $S_t P$ allenfalls eine meßbare $[P, t]$ -Funktion, und es bleibt dem alten äquivalent (vgl. D_2'), wir erhalten also eine äquivalente allgemeine Strömung. Der zulässige Änderungsbereich läßt sich auch so kennzeichnen: S_t innerhalb von $\bar{M}(x)$, falls $[x, t]$ zu $\bar{Y}_4$ gehört.

Wenn $\bar{N}$ meßbar und $\subseteq \Omega_1 \times T$ ist, so gilt

$$\bar{\mu}(\bar{N}) = \int_0^1 \bar{\mu}_{(x)}((M(x) \times T) \cdot \bar{N}) d\omega(x)$$

($\bar{\mu}_{(x)}$ das nach D_1 von $M(x)$ auf $M(x) \times T$ erweiterte $\mu_{(x)}$, der Integrand ist, bis auf eine x -Menge vom Maße 0 sinnvoll, d. h. $(M(x) \times T) \cdot \bar{N}$ $\bar{\mu}_{(x)}$ -meßbar). Dies folgt für $\bar{N} = O \times \tilde{O}$ (vgl. D_1) aus unserer ursprünglichen μ -, $\mu_{(x)}$ -Formel und wird auf die übrigen $\bar{N}$ so ausgedehnt, wie es für die analogen Formeln am Ende von 6 geschah.

Sei nun N meßbar, $\subseteq \mathcal{Q}_1$, und $\bar{N}$ sein $[P, t] \rightarrow S_t P$ -Urbild. Bis auf eine x -Menge vom Maße 0, $Y_5(N)$, ist $(M(x) \times T) \cdot \bar{N}$ ($\bar{\mu}_{(x)}$ -) meßbar, d. i. das $[P, t] \rightarrow [S_t P]$ -Urbild von $M(x) \cdot N$. Sei $Y_5 = \sum_{\xi \in A'} Y_5(K(\xi))$, auch dieses hat das Maß 0. Liegt y nicht in Y_5 , so sind die $[P, t] \rightarrow S_t P$ -Urbilder aller $M(y) \cdot K(\eta)$ ($\bar{\mu}_{(y)}$ -) meßbar, d. h. die Urbilder jener $M \subseteq M(y)$, die vermöge der $P \rightarrow [x, \xi]$ -Abbildung von 3. den Intervallen $x = y$, $0 \leq \xi \leq \eta$ entsprechen. Dasselbe gilt dann auch für diejenigen M , die Differenzen solcher Intervalle entsprechen, oder Summen endlich vieler solcher Differenzen, oder Limites aufsteigender Folgen solcher Summen — d. h. allen offenen Mengen auf $x = y$. Somit ist $S_t P$ als $[P, t]$ -Funktion in $M(x) \times T$ ($\bar{\mu}_{(x)}$ -) meßbar.

Für jedes x möge die Menge der t mit $[x, t]$ in $\bar{Y}_4$ $T(x)$ heißen. Nach dem Satz von Fubini ist $T(x)$ stets eine t -Menge vom Maße 0, abgesehen von einer x -Menge vom Maße 0, Y_4 .

Die Summe von Y_0, Y_4, Y_5 heiße Y_6 , dies ist auch eine Menge vom Maße 0, ihre Komplementärmenge (in $0 \leq x \leq 1$) heiße X_6 . Für die x von X_6 gilt folgendes: $\mu_{(x)}$ ist definiert, und ein l -Maß in $M(x)$. $T(x)$ hat das Maß 0; liegt t nicht in $T(x)$, so bildet $S_t M(x)$ minus eine Menge vom Maße 0 eineindeutig und maßtreu auf $M(x)$ minus eine Menge vom Maße 0 ab, für alle s bis auf eine Menge vom Maße 0 gilt $S_t S_s \cong S_{t+s}$ in $M(x)$ (wir können, ohne hieran zu ändern, auch $s, t+s$ als nicht in $T(t)$ liegend annehmen), $S_t P$ ist eine meßbare $[P, t]$ -Funktion (in $M(x)$ bzw. $M(x) \times T$, gemeint sind hier stets die Maße $\mu_{(x)}$ bzw. $\bar{\mu}_{(x)}$). Im nächsten Paragraphen werden wir zeigen, daß unter diesen Umständen die S_t der t in $T(x)$ so undefiniert bzw. neudefiniert werden können, daß diese S_t zusammen D'_2 in $M(x)$ (mit $\mu_{(x)}$) erfüllen, d. h. dort eine allgemeine Strömung bilden. Da für t von $T(x)$ $[x, t]$ zu $\bar{Y}_4$ gehört, bleibt nach dem früher Gesagten die neue S_t -Strömung der alten äquivalent. Wir dürfen also annehmen, daß das soeben Erreichte von vornherein der Fall war: daß für jedes x von X_6 die S_t eine allgemeine Strömung in $M(x)$ (mit $\mu_{(x)}$) bilden.

9. In 3. definierten wir: $\tilde{\mu}^{(0)}(X) = \int_X d\bar{\mu}(x, 1) = \int_X dw(x)$, und in 4. zeigten wir: $\tilde{\mu}^{(0)}(X) = \mu(M[X])$, was nach der μ -, $\mu_{(x)}$ -Formel gleich

$$\int_0^1 \mu_{(x)}(M(x) \cdot M[X]) dw(x) = \int_X \mu_{(x)}(M(x)) dw(x)$$

ist. Somit gilt allgemein

$$\int_X dw(x) = \int_X \mu_{(x)}(M(x)) dw(x),$$

es muß also, bis auf eine Menge Y_7 vom Maße 0, $\mu_{(x)}(M(x)) = 1$ sein.

Wir bilden zu jedem $M(x)$ den Hilbertschen Raum $\mathfrak{H}_{(x)}$ der in ihm definierten Funktionen $f_{(x)}(P)$ mit endlichem $\int_{M(x)} |f_{(x)}(P)|^2 dv_{(x,P)}$ (das Integral ist auf Grund des l -Maßes $\mu_{(x)}$ zu bilden, $v_{(x,P)}$ ist das entsprechende Volumelement). Da $P \rightarrow S_t P$ in $M(x)$ maßtreu ist, können wir die unitären Koopmanschen Operatoren $U_{(x,t)}$ bilden, und das dazugehörige $E_{(x,0)}$ (vgl. 1.).

Ferner sei $\chi_{(x,\xi)}(P) \begin{cases} = 1 & \text{in } M(x) \cdot K(\xi) \\ = 0 & \text{sonst in } M(x) \end{cases}$.

Ist $f_{(x)}(P)$ zu allen $\chi_{(x,\xi)}(P)$ (ξ in A') orthogonal, so wenden wir die Abbildung $P \rightarrow [y, \eta]$ aus 3. an. Dabei wird $y = x$, $0 \leq \eta \leq 1$. Aus $f_{(x)}(P)$ wird also ein $f'_{(x)}(\eta)$, aus $\chi_{(x,\xi)}(P)$ aber $\chi'_{(x,\xi)}(\eta) \begin{cases} = 1 & \text{für } 0 \leq \eta \leq \xi \\ = 0 & \text{sonst} \end{cases}$.

$f'_{(x)}$, $\chi'_{(x,\xi)}$ sind wieder orthogonal, d. h. $\int_0^\xi f'_{(x)}(\eta) d_{(x)} \eta = 0$ ($d_{(x)}$ deute an, daß das Integral auf Grund des $\tilde{\mu}_{(x)}$ -Maßes zu bilden ist) für alle ξ . Also ist $f'_{(x)}(\eta) = 0$, bis auf eine η -Menge vom ($\tilde{\mu}_{(x)}$ -) Maße 0, d. h. $f_{(x)}(P) = 0$, bis auf eine P -Menge vom ($\mu_{(x)}$ -) Maße 0. Da somit nur 0 zu allen $\chi_{(x,\xi)}$ (ξ in A') orthogonal ist, spannen diese die abgeschlossene Linearmannigfaltigkeit $\mathfrak{H}_{(x)}$ auf.

Wir setzen $\xi_0 = 1$, $\xi_1, \xi_2, \dots$ irgendeine Abzählung von A' , und orthogonalisieren die $\chi_{(x,\xi_n)}$, $n = 0, 1, 2, \dots$, in dieser Reihenfolge⁵³. So entstehen Funktionen $\varphi_{(x,m)} = \sum_{n=0}^{\infty} c_{(x,m,n)} \chi_{(x,\xi_n)}$, $m = 0, 1, 2, \dots$, die nach der obigen Bemerkung ein vollständiges normiertes Orthogonalsystem bilden. Da $\chi_{(x,\xi_0)} = 1$ (für alle x in $M(x)$) normiert ist, und der Prozeß mit ihm beginnt, ist $\varphi_{(x,0)} = \chi_{(x,\xi_0)} = 1$. Die $c_{(x,m,n)}$ entstehen durch elementare Rechenoperationen aus den

$$(\chi_{(x,\xi_m)}, \chi_{(x,\xi_n)}) = \mu_{(x)} (M(x) \cdot A(\xi_m) \cdot A(\xi_n)),$$

da diese meßbare x -Funktionen sind, sind es auch die $c_{(x,m,n)}$.

Sei t fest.

$$(S_t \chi_{(x,\xi_m)}, \chi_{(x,\xi_n)}) = \mu_{(x)} (M(x) \cdot S_t^{-1} A(\xi_m) \cdot A(\xi_n))$$

ist eine meßbare x -Funktion, da es die $c_{(x,m,n)}$ auch sind, sind es die $\varphi_{(x,m)}$ -Matrizenelemente der S_t , $s_{(t,x,m,n)} = (S_t \varphi_{(x,m)}, \varphi_{(x,n)})$ gleichfalls. Da nach 1

$$E_0 = \lim_{t \rightarrow \infty} \frac{1}{s-t} \int_t^s S_\tau d\tau = \lim_{p \rightarrow \infty} \left(\lim_{q \rightarrow \infty} \frac{1}{2pq} \sum_{r=-pq}^{+pq} S_{\frac{r}{q}} \right)$$

$$(p, q = 1, 2, \dots, \nu = 0, \pm 1, \pm 2, \dots)$$

⁵³ Nach dem Verfahren von E. Schmidt,

vgl. auch E S. 69.

ist⁵⁴, sind auch die $\varphi_{(x,m)}$ -Matrizenelemente der E_0 , $e_{(0,x,m,n)} = (E_0 \varphi_{(x,m)}, \varphi_{(x,n)})$ meßbare x -Funktionen.

Da $\varphi_{(x,0)}$ konstant ist, ist $E_0 \varphi_{(x,0)} = \varphi_{(x,0)}$, also da E_0 Projektionsoperator ist, $e_{(0,x,m,n)} \begin{cases} = 1, & \text{für } m = n = 0 \\ = 0, & \text{für } m = 0, n \neq 0 \\ = 0, & \text{für } m \neq 0, n = 0 \end{cases}$. Daß S_t in $M(x)$ ergodisch ist, besagt, daß $E_0 f_{(x)} = f_{(x)}$ keine Lösungen außer $f_{(x)} = c \varphi_{(x,0)}$ hat, d. h. $e_{(0,x,m,n)} = 0$, für $m, n \neq 0$. Die Menge der x , für die das nicht gilt, Y_8 , ist somit meßbar. Sei ein m gegeben, die x -Menge für die es bei diesem m und mindestens einem n nicht gilt, $Y_{8,n}$ ist ebenfalls meßbar, und es ist $Y_8 = \sum_{m=1}^{\infty} Y_{8,m}$. Wir werden zeigen, daß alle $Y_{8,m}$ das Maß 0 haben, also auch Y_8 . Nehmen wir also an, ein $Y_{8,m}$ hätte nicht das Maß 0.

Gehöre x zu $Y_{8,m}$. Da E_0 Projektionsoperator ist, ist $f_{(x)} = \sum_0^{\infty} e_{(0,x,m,n)} \varphi_{(x,n)}$ Lösung von $E_0 f_{(x)} = f_{(x)}$, also von $S_t f_{(x)} = f_{(x)}$, und nach dem was wir schon wissen, orthogonal zu $\varphi_{(x,0)} = \text{Konstanz}$ und $\neq 0$ (auch nicht bis auf Mengen vom Maße 0!). $f_{(x)}(P)$ ist aus den $e_{(0,x,m,n)}, c_{(x,m,n)}, \chi_{(x,\xi_m)}(P)$ aufgebaut. Lassen wir P über Ω_1 variieren, und x durch die Bedingung P in $M(x)$ bestimmen, so ist $\chi_{(x,\xi_m)}(P) \begin{cases} = 1, & \text{für } P \text{ in } K(\xi_m) \\ = 0, & \text{sonst} \end{cases}$, also (μ_-) -meß-

bar in P . $e_{(0,x,m,n)}, c_{(x,m,n)}$ sind $(\tilde{\mu}^{(0)})$ -meßbar in x , also auch (μ_-) -meßbar in P (denn jede meßbare x -Menge X hat als P -Urbild $M[X]$, das auch meßbar ist). Also ist $f_{(x)}(P)$ auch als P -Funktion (μ_-) -meßbar. Sei C die Menge der komplexen Zahlen $u + iv$ mit $u > 0$ oder $u = 0, v > 0$, D die mit $u < 0$ oder $u = 0, v < 0$. Dann sind auch die P -Mengen $f_{(x)}(P)$ in C und $f_{(x)}(P)$ in D (μ_-) -meßbar, also die x -Funktionen $\mu_{(x)}(P\text{-Menge in } M(x), f_{(x)}(P) \text{ in } C)$, $\mu_{(x)}(P\text{-Menge in } M(x), f_{(x)}(P) \text{ in } D)$ beide $(\tilde{\mu}^{(0)})$ -meßbar. Bei festem x ist $f_{(x)}(P)$ (P in $M(x)$) nicht bis auf eine Menge vom Maße 0 = 0, und zur Konstanten orthogonal, also muß es in Mengen vom Maße > 0 in C liegen und in D liegen. D. h. die obigen Funktionen sind stets > 0 . Wir können also ein $\epsilon > 0$ finden, so daß beide in einer Teilmenge Y' von Y_8 , die ein Maß > 0 hat, stets $\geq \epsilon$ sind. Wir können dabei Y' auch mit endlichem Maße annehmen. Wir nennen nun die Menge der P aus $M[Y']$ mit $f_{(x)}(P)$ in C bzw. D , J bzw. K .

Nach Entstehen sind J, K fremd, beide (μ_-) -meßbar. Die P -Menge, auf der sich $f_{(x)}(P), f_{(x)}(S_t P)$ unterscheiden, ist (μ_-) -meßbar, für jedes x ist ihr in $M(x)$ liegender Teil (wegen $U_t f_{(x)} = f_{(x)}$ in $M(x)$) vom $(\mu_{(x)})$ -Maße 0, also hat sie (nach unserer μ_- , $\mu_{(x)}$ -Formel) selbst das Maß 0. Also sind

⁵⁴ Die Limites gelten im Sinne der starken Operatoren-Konvergenz, und wir brauchen nur die Konvergenz der Matrizenelemente, d. h. die schwache. Vgl. a. a. O. Anm. ⁴³.

J, K bis auf Mengen vom (μ_-) Maße 0, S_t -invariant. Sie sind ferner $\subseteq M[Y']$, $\mu(M[Y']) = \tilde{\mu}^{(0)}(Y')$ ist endlich, also sind J, K $\mathcal{A}$ -Mengen. Nach dem am Ende von 5. Gesagten sind sie also, bis auf Mengen vom (μ_-) Maße 0, gleich $M[Z], M[W]$.

Da aber in Y' stets $\mu_{(x)}(M(x) \cdot J) \geq \varepsilon$ ist, gilt für jedes $(\tilde{\mu}^{(0)}$ -meßbare) $V \subseteq Y'$

$$\mu_{(x)}(M(V) \cdot J) \geq \varepsilon \cdot \tilde{\mu}^{(0)}(V) \quad (\mu_-, \mu_{(x)}\text{-Formel!}),$$

ebenso

$$\mu(M(V) \cdot K) \geq \varepsilon \cdot \tilde{\mu}^{(0)}(V).$$

Setzen wir in der ersten Formel $V = W$, so wird

$$\begin{aligned} \varepsilon \cdot \tilde{\mu}^{(0)}(W) &\leq \mu(M(W) \cdot J) = \mu(K \cdot J) = 0, \\ \tilde{\mu}^{(0)}(W) &= 0, \quad \mu(K) = \mu(M[W]) = 0, \end{aligned}$$

und setzen wir nun in der zweiten Formel $V = Y'$, so wird

$$\varepsilon \cdot \tilde{\mu}^{(0)}(Y') \leq \mu(M[Y'] \cdot K) = \mu(K) = 0, \quad \tilde{\mu}^{(0)}(Y') = 0.$$

Aber dies widerspricht der ursprünglichen Annahme über Y' .

Somit muß Y_8 das Maß Null haben.

10. Sei Y_9 die Summe von Y_6, Y_7, Y_8 , also auch vom Maße 0, X_9 seine Komplementärmenge (in $0 \leq x \leq 1$). Für die x von X_9 ist also $\mu_{(x)}$ ein l -Maß in $M(x)$, S_t ein allgemeiner Fluß in $M(x)$, $\mu_{(x)}(M(x)) = 1$, S_t ergodisch in $M(x)$.

Es ist $\mu(M[Y_9]) = \tilde{\mu}^{(0)}(Y_9) = 0$, wenn wir es also zu Ω_2 schlagen, behält dieses seine Eigenschaften aus 3., und es wird $\mu(M[X_9]) = \Omega_1$. Hat dieses Ω_2 das Maß 0, so können wir es zu einem $M(x)$ schlagen, $\mu_{(x)}(\Omega_2) = 0$ definierend, dann bleibt $M[X_9] = \Omega_1$ erhalten, und es wird $\Omega_2 = 0$. Wir können also annehmen, daß diese Forderungen erfüllt sind, ferner wollen wir X_9 lieber $X^{(0)}$ nennen. Dann können wir zusammenfassend aussagen:

SATZ 2. Ω kann folgendermaßen zerlegt werden: $\Omega = \Omega_1 + \Omega_2$, $\Omega_1 = \sum_{x \in X^{(0)}} M(x)$,

wobei $X^{(0)}$ eine Menge von Zahlen in $0 \leq x \leq 1$ ist, und die $M(x)$ sowie Ω_2 paarweise elementfremde Mengen. $X^{(0)}$ und die $M(x)$ sind m -Räume, in ihnen sind l -Maße $\tilde{\mu}^{(0)}$ bzw. $\mu_{(x)}$ definiert. Diese Zerlegung hat, nachdem evtl. S_t durch eine äquivalente allgemeine Strömung ersetzt wurde, die folgenden Eigenschaften:

1. Ω_2 ist entweder leer, oder $\mu(\Omega_2) = \infty$ und Ω_2 null-ergodisch.
2. In jedem $M(x)$ ist S_t (unter Zugrundelegung des l -Maßes $\mu_{(x)}$ und der auf $M(x)$ übertragenen relativen Topologie aus Ω) eine eigentlich-ergodische allgemeine Strömung.

3. Es ist für alle $(\tilde{\mu}^{(0)})$ -meßbaren $X \subseteq X^{(0)}$ $\tilde{\mu}^{(0)}(X) = \mu\left(\sum_{x \text{ in } X} M(x)\right)$. Es ist stets $\mu_{(x)}(M(x)) = 1$.

Es ist für alle (μ) -meßbaren $N \subseteq \Omega_1$ $\mu(N) = \int_X \mu_{(x)}(M(x) \cdot N) d w(x)$.

D. h.: der Integrand ist, bis auf eine x -Menge vom $(\tilde{\mu}^{(0)})$ -Maße 0, sinnvoll, also $M(x) \cdot N$ $(\mu_{(x)})$ -meßbar; $\mu_{(x)}(M(x) \cdot N)$ ist als x -Funktion $(\tilde{\mu}^{(0)})$ -meßbar; $d w(x)$ ist das Differentialelement der $\tilde{\mu}^{(0)}$ -Integration, also

$$w(y) = \tilde{\mu}^{(0)}(\text{Menge } 0 \leq x \leq y) = \mu\left(\sum_{0 \leq x \leq y} M(x)\right);$$

und außerdem wird die Gleichheit beider Seiten behauptet.

Damit haben wir bewiesen, was in Einleitung 4. angekündigt wurde: daß jede allgemeine Strömung in ergodische Strömungen zerlegt werden kann. Es ließe sich noch zeigen, daß die Zerlegung von Satz 2 „im Wesentlichen“ nur auf eine Weise möglich ist, wir gehen aber darauf nicht mehr ein.

III. Charakteristische Eigenschaften der Koopmanschen U_t -Scharen.

1. Betrachten wir zunächst eine einzige eindeutige maßtreue Abbildung $P \rightarrow SP$ von Ω' auf Ω'' , oder auch nur von Ω' minus einer Menge vom Maße 0 auf Ω'' minus eine Menge vom Maße 0. Daß der durch $Uf(P) = f(SP)$ definierte Operator U unitär ist, ist klar, ebenso daß $(D_s$ entsprechend) $Uf \cdot Ug = U(f \cdot g)$ ist, falls $f, g, f \cdot g$ zum Hilbertschen Raume $\mathfrak{H}$ gehören. Letzteres ist z. B. bestimmt der Fall, wenn f, g zum Hilbertschen Raume gehören und beschränkt sind⁵⁵. Hat nun umgekehrt ein unitärer Operator U in $\mathfrak{H}$ diese Eigenschaft: $Uf \cdot Ug = U(f \cdot g)$ für beschränkte f, g , so können wir ein eindeutiges und maßtreues S finden, das U nach $Uf(P) = f(SP)$ erzeugt.

Sei nämlich $\mathcal{A}' \subseteq \Omega$ meßbar und von endlichem Maße, dann ist für $\chi_{\mathcal{A}'}(P) \Big| = 1 \text{ für } P \text{ in } \mathcal{A}' \Big| \chi_{\mathcal{A}'} \cdot \chi_{\mathcal{A}'} = \chi_{\mathcal{A}'}$, also auch $U\chi_{\mathcal{A}'} \cdot U\chi_{\mathcal{A}'} = U(\chi_{\mathcal{A}'} \cdot \chi_{\mathcal{A}'}) = U\chi_{\mathcal{A}'}$. Somit ist, bis auf eine Menge vom Maße 0, $U\chi_{\mathcal{A}'}(P) = 0$ oder 1, und da wir es auf dieser etwa zu 0 umdefinieren können, ohne es als $\mathfrak{H}$ -Element zu ändern, dürfen wir annehmen, daß $U\chi_{\mathcal{A}'}(P) = 0$ oder 1 stets gilt. Die Menge, wo letzteres gilt, heiße $\mathcal{A}''$, sie ist meßbar (weil $U\chi_{\mathcal{A}'}$ es ist) und von endlichem Maße (wie $U\chi_{\mathcal{A}'}$ zu $\mathfrak{H}$ gehört). Es ist somit $U\chi_{\mathcal{A}'} = \chi_{\mathcal{A}''}$.

⁵⁵ Denn es ist

$$\int_{\Omega} |f(P)g(P)|^2 dv_P \leq \left(\max_{P \text{ in } \Omega} |f(P)|\right)^2 \int_{\Omega} |g(P)|^2 dv_P,$$

also endlich.

Wir haben durch diese Zuordnung jedem meßbaren $\mathcal{A}' \subseteq \Omega'$ von endlichem Maße ein ebensolches $\mathcal{A}'' \subseteq \Omega''$ zugeordnet; da U^{-1} ebenso ist, gilt auch die Umkehrung, d. h. jedes solche $\mathcal{A}''$ tritt bei einem gewissen $\mathcal{A}'$ wirklich auf. Ferner legt $\mathcal{A}'$ offenbar $\mathcal{A}''$ bis auf eine Menge vom Maße 0 fest, ebenso $\mathcal{A}''$ das $\mathcal{A}'$.

Da U unitär ist und $\|\chi_{\mathcal{A}'}\|^2 = \mu(\mathcal{A}')$, $\|\chi_{\mathcal{A}''}\|^2 = \mu(\mathcal{A}'')$, ist $\mu(\mathcal{A}') = \mu(\mathcal{A}'')$. Es ist

$$U\chi_{\mathcal{A}_1 \cdot \mathcal{A}_2} = U(\chi_{\mathcal{A}_1} \cdot \chi_{\mathcal{A}_2}) = U\chi_{\mathcal{A}_1} \cdot U\chi_{\mathcal{A}_2} = \chi_{\mathcal{A}_1'} \cdot \chi_{\mathcal{A}_2'} = \chi_{\mathcal{A}_1' \cdot \mathcal{A}_2'}$$

und

$$\begin{aligned} U\chi_{\mathcal{A}_1 + \mathcal{A}_2} &= U(\chi_{\mathcal{A}_1} + \chi_{\mathcal{A}_2} - \chi_{\mathcal{A}_1} \cdot \chi_{\mathcal{A}_2}) = U\chi_{\mathcal{A}_1} + U\chi_{\mathcal{A}_2} - U\chi_{\mathcal{A}_1 \cdot \mathcal{A}_2} \\ &= \chi_{\mathcal{A}_1'} + \chi_{\mathcal{A}_2'} - \chi_{\mathcal{A}_1' \cdot \mathcal{A}_2'} = \chi_{\mathcal{A}_1' + \mathcal{A}_2'}. \end{aligned}$$

D. h.: wenn $\mathcal{A}_1, \mathcal{A}_2$ bzw. $\mathcal{A}_1', \mathcal{A}_2'$ entsprechen, entspricht $\mathcal{A}_1 \cdot \mathcal{A}_2$ $\mathcal{A}_1' \cdot \mathcal{A}_2'$ und $\mathcal{A}_1 + \mathcal{A}_2$ $\mathcal{A}_1' + \mathcal{A}_2'$. Durch Induktion folgt: wenn $\mathcal{A}_1, \dots, \mathcal{A}_n$ bzw. $\mathcal{A}_1', \dots, \mathcal{A}_n'$ entsprechen, so entspricht $\mathcal{A}_1 + \dots + \mathcal{A}_n$ bzw. $\mathcal{A}_1' + \dots + \mathcal{A}_n'$. Und da U ein stetiger Operator ist, ergibt ein Grenzübergang: wenn $\mathcal{A}_1, \mathcal{A}_2, \dots$ bzw. $\mathcal{A}_1', \mathcal{A}_2', \dots$ entsprechen (und $\mathcal{A}_1 + \mathcal{A}_2 + \dots$ ein endliches Maß hat), so entspricht $\mathcal{A}_1 + \mathcal{A}_2 + \dots$, $\mathcal{A}_1' + \mathcal{A}_2' + \dots$.

Wir haben somit in $\mathcal{A}' \leftrightarrow \mathcal{A}''$, im Sinne von Satz 1 der vorhergehenden Abhandlung des Verfassers, eine maßtreue Mengenabbildung von Ω' auf Ω'' vor uns, und ebenso ist $\mathcal{A}'' \leftrightarrow \mathcal{A}'$. Dann existiert aber auf Grund des genannten Satzes eine eindeutige maßtreue Punktabbildung $P \rightarrow SP$ von Ω' minus einer Menge vom Maße 0 auf Ω'' minus eine Menge vom Maße 0, die diese Mengenabbildung $\mathcal{A}'' \leftrightarrow \mathcal{A}'$ erzeugt. Wenn Ω', Ω'' un abzählbar sind, können wir S sogar Ω' eindeutig auf Ω'' abbildend wählen.

$Vf(P) = f(SP)$ ist also unitär. Es ist $U\chi_{\mathcal{A}'} = \chi_{\mathcal{A}''}$, aber $\mathcal{A}''$ ist das $P \rightarrow SP$ -Urbild von $\mathcal{A}'$, also auch $V\chi_{\mathcal{A}'} = \chi_{\mathcal{A}''}$. Somit gilt $Uf = Vf$ für alle $f = \chi_{\mathcal{A}'}$. Da aber diese die abgeschlossene Linearmannigfaltigkeit $\mathfrak{H}$ aufspannen (vgl. a. a. O. Anm. ⁴⁵), gilt es für alle f . Also ist $U = V$, womit alles bewiesen ist.

2. Sei nun S_t eine allgemeine Strömung, U_t die Koopmanschen Operatoren: $U_t f(P) = f(S_t P)$. Unitarität: $U_t U_s = U_{t+s}$, $U_t f \cdot U_t g = U_t (f \cdot g)$ (f, g beschränkt) sind klar, es fehlt nur noch eins aus D_4, D_5 , daß $(U_t f, g)$ eine meßbare Funktion von t ist. Dies genügt es aber, für die f, g aus einer solchen Menge $\mathfrak{R}$ zu beweisen, die die abgeschlossene Linearmannigfaltigkeit $\mathfrak{H}$ aufspannt, denn wenn $(U_t f, g)$ für die f, g von $\mathfrak{R}$ meßbar ist, ist es es auch für deren Linearaggregate und die Limites dieser Linearaggregate, d. h. für alle f, g . Eine solche Menge $\mathfrak{R}$ ist nun die Menge aller $\chi_{\mathcal{A}}$ ($\mathcal{A}$ meßbar, von endlichem Maß, vgl. a. a. O. Anm. ⁴⁵), wir haben also nur zu zeigen: $(U_t \chi_{\mathcal{A}}, \chi_{\mathcal{M}}) = (\chi_{S_t^{-1} \mathcal{A}}, \chi_{\mathcal{M}}) = \mu(S_t^{-1} \mathcal{A} \cdot \mathcal{M})$ ist eine

meßbare Funktion von t . In II. 6. zeigten wir, daß das $[P, t] \rightarrow S_t P$ -Urbild von $\mathcal{A}$ (in $\Omega \times T$) meßbar ist, und die Menge der $[P, t]$, P in M , ist offenbar auch meßbar. Folglich ist der Durchschnitt dieser zwei Mengen, N , meßbar. Nach dem Satz von Fubini ist also $\mu(N_{(t)})$ ($N_{(t)}$ ist die P -Menge mit $[P, t]$ in N) eine meßbare Funktion von t . Da aber $N_{(t)} = S_t^{-1} \mathcal{A} \cdot M$ ist, sind wir damit am Ziel.

Schlagen wir nun den umgekehrten Weg ein: sei U_t eine D_4, D_5 erfüllende Operatorenschar. Zunächst können wir jedem U_t ein $P \rightarrow S_t P$ zuordnen, das Ω minus eine Menge vom Maße 0 eineindeutig und maßtreu auf Ω minus eine Menge vom Maße 0 abbildet. Aus $U_t U_s = U_{t+s}$ folgt $S_t S_s = S_{t+s}$. Es bleibt noch zu prüfen, ob $S_t P$ meßbar von $[P, t]$ abhängt.

Aus Satz 1 der vorhergehenden Abhandlung des Verfassers folgt, daß U_t stark stetig in t ist, also insbesondere $U_t f \rightarrow f$ für $t \rightarrow 0$. Ist $f = \chi_A$, so besagt dies $\chi_{S_t^{-1}A} \rightarrow \chi_A$, und da $\|\chi_{S_t^{-1}A} - \chi_A\|^2 = \mu((\mathcal{A} + S_t^{-1}\mathcal{A}) - (\mathcal{A} \cdot S_t^{-1}\mathcal{A}))$ ist, $\mu((\mathcal{A} + S_t^{-1}\mathcal{A}) - (\mathcal{A} \cdot S_t^{-1}\mathcal{A})) \rightarrow 0$. Wenn von $P, S_t P$ beide oder keins in $\mathcal{A}$ liegen, sagen wir, sie liegen in bezug auf $\mathcal{A}$ gleich; wie wir sehen, strebt das Maß der Menge der mit ihrem S_t -Bilde in bezug auf $\mathcal{A}$ nicht gleich liegenden P mit $t \rightarrow 0$ gegen Null. Sei nun $K_1, K_2, \dots$ ein „gleichwertiges Umgebungssystem“ in Ω . Sei $\epsilon > 0$ und ein $n = 1, 2, \dots$ gegeben. Das Maß der Menge der mit ihrem S_t -Bilde in bezug auf K_m nicht gleich liegenden P strebt mit $t \rightarrow 0$ gegen Null, also auch das Maß der Summe dieser Mengen für alle $m = 1, \dots, n$. Es gibt also ein $\delta = \delta(n, \epsilon) > 0$, so daß dieses Maß $\leq \epsilon$ ist, wenn $|t| \leq \delta$ ist. Sei nun $|t' - t| \leq \delta$. Dann hat die Menge der P , die mit ihrem $S_{t'-t}$ -Bilde nicht in bezug auf alle $K_1, \dots, K_n$ gleich liegen, ein Maß $\leq \epsilon$, also auch ihr S_t -Urbild, d. i. die Menge $\Delta(t', t'', n)$ aller P , für welche das $S_{t'}$ - und das $S_{t''}$ -Bild nicht in bezug auf alle $K_1, \dots, K_n$ miteinander gleich liegen.

Sei $\delta_n = \delta\left(n, \frac{1}{2^n}\right)$ und eine Folge $t_1, t_2, \dots$ sowie ein t gegeben mit

$|t - t_n| \leq \delta_n$. Die Menge $\Delta_n = \sum_{m=n}^{\infty} (t_m, t, m)$ hat ein Maß $\leq \sum_{m=n}^{\infty} \frac{1}{2^m} = \frac{1}{2^{n-1}}$.

Liegt P nicht in ihr, so liegt $S_{t_n} P$ mit $S_t P$ in bezug auf alle $K_1, \dots, K_n$ gleich. Da dann P erst recht in keinem Δ_p , $p \geq n$ liegt, liegen die $S_{t_p} P$ mit $S_t P$ in bezug auf alle $K_1, \dots, K_p$ gleich, also erst recht in bezug auf alle $K_1, \dots, K_n$. D. h.: alle $S_{t_p} P$, $p \geq n$ sowie $S_t P$ liegen miteinander in bezug auf alle $K_1, \dots, K_n$ gleich.

Liege nun P nicht in Δ_n , sei K eine Umgebung von $S_t P$. Dann gibt es ein $S_t P$ enthaltendes $K_m \subseteq K$. Sei $p \geq m, n$, dann liegen nach Obigem alle $S_{t_q} P$, $q \geq p$, mit $S_t P$ in bezug auf alle $K_1, \dots, K_p$ gleich (denn P gehört erst recht nicht zu Δ_p), also auch in bezug auf K_m . Da $S_t P$ in K_m liegt, liegen somit auch die $S_{t_q} P$ dort, also erst recht in K . D. h.: alle

$S_t P$ mit $q \geq p$ liegen in K . Da dies für jedes K gilt, haben wir $S_t P \rightarrow S_t P$ für $q \rightarrow \infty$. Die Menge der P , für die nicht $S_t P \rightarrow S_t P$ für $q \rightarrow \infty$ ist, ist also jedenfalls $\subseteq \Delta_n$. Ihr (äußeres) Maß ist somit $\leq \frac{1}{2^{n-1}}$. Und da dies für jedes n gilt, hat sie das Maß 0.

Wir haben gezeigt: wenn $|t - t_n| \leq \delta_n$ ist, so gilt, bis auf eine P -Menge vom Maße 0, $S_{t_n} P \rightarrow S_t P$ für $n \rightarrow \infty$. Nun definieren wir für jedes $n = 1, 2, \dots$: $\Sigma_t^{(n)} P = S_{\nu \delta_n} P$, wobei $\nu = \nu(n, t)$ aus $\nu \delta_n \leq t < (\nu + 1) \delta_n$ zu bestimmen ist. In einem Intervalle $\nu_0 \delta_n \leq t < (\nu_0 + 1) \delta_n$ ist $\Sigma_t^{(n)} P$ von t unabhängig, und in seiner P -Abhängigkeit ein $S_{t_0} P$ ($t_0 = \nu_0 \delta$), also sowohl P - als auch $[P, t]$ -meßbar. Da aber eine Folge solcher Intervalle den ganzen $[P, t]$ -Raum ($\Omega \times T$) ausmacht, ist es überhaupt $[P, t]$ -meßbar. Wir können nun ein $S'_t P$ so definieren: wenn $\lim_{n \rightarrow \infty} \Sigma_t^{(n)} P$ existiert, sei $S'_t P$ dieser Limes,

sonst sei es undefiniert. Dieses $S'_t P$ ist offenbar auch meßbar. Da nun $|t - \nu \delta_n| \leq \delta_n$ ist, folgt aus dem weiter oben Bewiesenen (mit $t_n = \nu(n, t) \delta_n$), daß bei festem t $S'_t P$ bis auf eine P -Menge vom Maße 0 definiert und gleich $S_t P$ ist. Somit bilden auch die $S'_t \Omega$ minus eine Nullmenge eineindeutig und maßtreu auf Ω minus eine Nullmenge ab, auch sie erzeugen die unitären Operatoren U_t , auch für sie gilt $S'_t S'_s = S'_{t+s}$. Damit D'_2 erfüllt sei, ist nur noch eines notwendig: $S'_t P$ in jenen P -Mengen vom Maße 0 undefiniert machen, die bei der eineindeutigen Abbildung $P \rightarrow S_t P$ zu streichen sind — und dies unter Wahrung seiner $[P, t]$ -Meßbarkeit.

Daß das $[P, t] \rightarrow [S'_t P, t]$ -Urbild jeder meßbaren Menge $\bar{N}$ meßbar ist, zeigt man wie in 6., also ist es auch das $[P, t] \rightarrow S'_t P$ -Urbild jeder meßbaren Menge N (man setze $\bar{N} = N \times T$). Daher ist bei festem s $S'_s S'_t P$ eine meßbare $[P, t]$ -Funktion. Da $\Sigma_t^{(n)} S'_t P$ aus abzählbar vielen solchen zusammengesetzt ist, ist es ebenfalls $[P, t]$ -meßbar, und sein Limes für $n \rightarrow \infty$, $S'_{-t} S'_t P$ ist es auch. Die $[P, t]$ -Menge $\bar{M}$ mit $S'_{-t} S'_t P \neq P$ ist also meßbar, und für jedes t hat die Menge der P mit $[P, t]$ in $\bar{M}$ das Maß 0. Nach dem Satze von Fubini hat also auch $\bar{M}$ das Maß 0. Erklären wir daher $S'_t P$ in $\bar{M}$ für undefiniert, so sind wir am Ziele.

Wir haben bewiesen:

SATZ 3. *Damit eine Operatorenschar nach Einleitung 3. zu einer allgemeinen Strömung S_t gehöre, d. h. damit sie eine Koopmansche Schar sei, ist notwendig und hinreichend, daß sie D_4, D_5 (Einleitung 3.) erfüllt. (D_5 braucht bloß für beschränkte f, g gefordert zu werden.)*

3. Wir holen jetzt noch den Beweis der in II., 8. verwendeten Aussage nach. Es sei eine Schar eineindeutiger und maßtreuer Abbildungen S_t von Ω minus einer Menge vom Maße 0 auf Ω minus eine Menge vom Maße 0 gegeben, jedoch sei S_t nur für die t außerhalb einer gewissen t -Menge T'

vom (gewöhnlichen Lebesgueschen) Maße 0 definiert, und $S_t S_s \cong S_{t+s}$ braucht auch bei festem t nur bis auf eine s -Menge vom Maße 0 zu gelten, schließlich sei $S_t P$ als $[P, t]$ -Funktion meßbar. Behauptet wird: es ist möglich, die S_t für die t von T' so hinzu zu definieren, daß ein allgemeiner Fluß entsteht (d. h. D'_2 erfüllt ist).

Zu den gegebenen S_t (t nicht in T') können wir jedenfalls die U_t bilden, sie sind unitär, und für jedes gegebene t (nicht in T') gilt für alle s , bis auf eine s -Menge vom Maße 0, $U_t U_s = U_{t+s}$. Ferner ist $(U_t f, g)$ meßbar in t , was genau so bewiesen wird, wie am Anfang von 2. Könnten wir nun für die t von T' die U_t so hinzudefinieren, daß ausnahmslos $U_t U_s = U_{t+s}$ gilt, so erfüllen die neuen U_t D_4, D_5 , denn auch die t -Meßbarkeit von $(U_t f, g)$ bleibt erhalten, da die Neudefinition in einer t -Menge vom Maße 0 erfolgte. Nach Satz 3 gibt es dann eine allgemeine Strömung S'_t , die diese U_t erzeugt. Für die t außerhalb von T' erzeugen S_t, S'_t dasselbe U_t , es ist also $S_t \cong S'_t$. Wenn wir also S'_t außerhalb von T' in S_t undefinieren, entsteht eine äquivalente allgemeine Strömung, S'_t , und diese ist die gesuchte Erweiterung von S_t .

Die Aufgabe ist also: T' hat das Maß 0; für jedes t außerhalb T' ist ein unitäres U_t gegeben, und es ist für alle s , bis auf eine (von t abhängige) s -Menge vom Maße 0, $U_t U_s = U_{t+s}$; $(U_t f, g)$ ist eine meßbare t -Funktion — für die t von T' soll U_t so hinzudefiniert werden, daß ausnahmslos $U_t U_s = U_{t+s}$ gilt.

Da stets $1 \cdot U_s = U_{0+s}$ ist, ist U_0 , falls 0 nicht zu T' gehört, gleich 1; und wenn 0 zu T' gehört, können wir es aus T' herausnehmen und $U_0 = 1$ definieren. 0 nicht in T' , $U_0 = 1$, darf also jedenfalls vorausgesetzt werden.

Bei gegebenem t bilden die r von T' , sowie die r mit $t+r$ in T' je eine Menge vom Maße 0, wir können also r so wählen, daß weder r , noch $s = t+r$ zu T' gehört. D. h.: s, r nicht in T' , $t = s - r$, kann erreicht werden. Wir setzen nun $U'_t = U_s U_r^{-1}$, wobei freilich noch zu zeigen ist, daß U'_t von der Wahl von s, r (innerhalb der obigen Bedingungen) nicht abhängt. Ist aber dies gezeigt, so schließen wir so weiter: Für t nicht in T' können wir $r, t+r$ außerhalb von T' wählen, so daß $U_t U_r = U_{t+r}$ gilt, dann ist $U_t = U_{t+r} U_r^{-1} = U'_t$. Für beliebige t, s können wir $r, s+r, t+s+r$ außerhalb von T' wählen, dann ist

$$U'_t U'_s = U_{t+s+r} U_{s+r}^{-1} U_{s+r} U_r^{-1} = U_{t+s+r} U_r^{-1} = U'_{t+s}.$$

D. h. U'_t ist die gesuchte Definitionserweiterung für U_t . Es ist also nur noch zu zeigen: aus s, r, s', r' nicht in T' , $s - r = s' - r'$ folgt $U_s U_r^{-1} = U_{s'} U_{r'}^{-1}$.

Es genügt zu zeigen: aus s, r, s', r' nicht in T' , $s+r = s'+r'$ folgt $U_s U_r = U_{s'} U_{r'}$. Denn mit $r' = s, s' = r$ ergibt das $U_s U_r = U_r U_s$, d. h. die U_t kommutieren miteinander. Und damit wird $U_s U_r = U_{s'} U_{r'}$ $U_s U_{r'}^{-1} = U_{s'} U_r^{-1}$ gleichbedeutend. $s+r = s'+r'$ ist $s-r' = s'-r$ gleichbedeutend, wir haben also die ursprüngliche Behauptung wieder, bloß mit s, r', s', r an Stelle von s, r, s', r' .

Um die obige Behauptung zu beweisen, betrachten wir den Beweis von Satz 1 in der zweitvorhergehenden Abhandlung des Verfassers. Der Beweis dafür, daß für alle t außerhalb einer Menge $T'' (\supseteq T')$ vom Maße 0 U_t approximativ stark-stetig ist (vgl. a. a. O.), gilt unverändert. Gehört t nicht zu T' , so ist für alle s , bis auf eine Menge vom Maße 0, $U_t U_s = U_{t+s}$, also

$$\begin{aligned} \|U_{t+s}f - U_t f\| &= \|U_t U_s f - U_t U_0 f\| = \|U_t(U_s f - U_0 f)\| \\ &= \|U_s f - U_0 f\|. \end{aligned}$$

Wenden wir dies auf ein t außerhalb von T'' an, so ergibt es, daß U_t insbesondere für $t=0$ approximativ stark-stetig ist; wenden wir es auf die t von T' an, so zeigt sich, daß aus der approximativen starken Stetigkeit von U_t für $t=0$ dieselbe für alle t von T' folgt — d. h. $T' = T''$.

Seien nun s, r, s', r' gegeben, nicht in T' , $s+r = s'+r'$. Sei f aus $\mathfrak{S}$, $\varepsilon > 0$. Dann gibt es (da U_t für $t=r$ und $t=r'$ approximativ stark-stetig ist) ein $\delta_1 > 0$, so daß für jedes $\delta \leq \delta_1$, > 0 in $|t-r| \leq \delta$

$\|U_t f U_r f\| \leq \frac{\varepsilon}{2}$ gilt, bis auf eine Menge vom Maße $< \delta$, ebenso ein

$\delta_2 > 0$, so daß für jedes $\delta \leq \delta_2$, > 0 in $|t'-r'| \leq \delta$ $\|U_{t'} f - U_{r'} f\| \leq \frac{\varepsilon}{2}$

gilt, bis auf eine Menge vom Maße $< \delta$. Ferner gilt bis auf Mengen vom Maße 0 $U_s U_t = U_{s+t}$, sowie $U_{s'} U_{t'} = U_{s'+t'}$. Sei $\delta = \min(\delta_1, \delta_2)$. Binden wir also t, t' durch die Bedingung $s+t = s'+t'$ aneinander, so gelten, bis auf eine Menge vom Maße $< 2\delta$ alle vier Bedingungen. In $|t-r| < \delta$, das $|t'-r'| < \delta$ gleichbedeutend ist, und das Maß 2δ hat, gibt es also ein t , das sie alle erfüllt. Also ist

$$U_s U_t = U_{s+t} = U_{s'+t'} = U_{s'} U_{t'},$$

$$\|U_s U_t f - U_s U_r f\| = \|U_s (U_t f - U_r f)\| = \|U_t f - U_r f\| \leq \frac{\varepsilon}{2},$$

$$\|U_{s'} U_{t'} f - U_{s'} U_{r'} f\| = \|U_{s'} (U_{t'} f - U_{r'} f)\| = \|U_{t'} f - U_{r'} f\| \leq \frac{\varepsilon}{2},$$

also $\|U_s U_r f - U_{s'} U_{r'} f\| \leq \varepsilon$. Da dies für jedes $\varepsilon > 0$ gilt, ist $U_s U_r f = U_{s'} U_{r'} f$, und da dies für jedes f gilt, $U_s U_r = U_{s'} U_{r'}$, wie behauptet wurde.

Damit ist der notwendige Beweis nachgeholt.

IV. Die U_t -Scharen mit Punktspektrum.

1. In Einleitung 3 haben wir den auf dem Satz von Stone beruhenden Zusammenhang zwischen der U_t -Schar, der Zerlegung der Einheit $E(\lambda)$, und dem hypermaximalen Operator A angegeben. Es ist bekannt, was unter dem Spektrum von A zu verstehen ist, und wie dieses eingeteilt wird⁵⁷, wir wollen es auch Spektrum der $E(\lambda)$ oder Spektrum der U_t nennen, da es diese auch festlegen. Im Folgenden sollen die Eigenschaften dieses Spektrums untersucht werden.

Auf Grund von Satz 2 in § II brauchen wir nur solche U_t -Scharen zu untersuchen, die zu ergodischen Strömungen gehören. D. h. für die bei $\lim_{\substack{\varepsilon > 0 \\ \varepsilon \rightarrow 0}} (E(+\varepsilon) - E(-\varepsilon)) = \lim_{\substack{\varepsilon > 0 \\ \varepsilon \rightarrow 0}} (E(0) - E(-\varepsilon))$ ($E(\lambda)$ ist nach rechts

halbstetig!) $= E_0$ nur für konstantes f $E_0 f = f$ gilt. Bei nullergodischer Strömung ($\mu(\Omega) = \infty$) gehört außer 0 kein konstantes f zu $\mathfrak{E}$, also ist $E_0 = 0$; bei eigentlich-ergodischer Strömung ($\mu(\Omega)$ endlich) gehören die Konstanten f zu $\mathfrak{E}$, sie sind die $c\varphi_0$, wo φ_0 das normierte $\frac{1}{V\mu(\Omega)}$ ist,

also ist $E_0 = P_{\{\varphi_0\}}$ ⁵⁸. Im ersten Falle gehört also 0 nicht zum Punktspektrum; im zweiten gehört es dazu, u. zw. als einfacher Eigenwert mit der Eigenfunktion $\varphi_0 = \frac{1}{V\mu(\Omega)}$.

Der nullergodische Fall wurde von Koopman diskutiert (a. a. O. Anm. ²). Gäbe es dort ein Punktspektrum, so wäre für ein $f \neq 0$ $Af = \lambda f$, d. h. $U_t f = e^{i\lambda t} f$. Hieraus folgt

$$|U_t f(P)|^{59} = |U_t f(P)| = |e^{i\lambda t} \cdot f(P)| = |f(P)|.$$

Also ist $|f(P)|$ Eigenfunktion vom Eigenwerte 0, da aber 0 kein Eigenwert ist, $|f(P)| = 0$, $f(P) = 0$. Es gibt also kein Punktspektrum, sondern nur ein reines Streckenspektrum.

Im eigentlich-ergodischen Falle können wir $\mu(\Omega) = 1$ annehmen (vgl. Satz 2). Wieder folgt nach Koopman für jede Eigenfunktion f des Eigenwertes λ , daß $|f|$ Eigenfunktion des Eigenwertes 0 ist. Hier bedeutet dies aber nur, daß $|f|$ konstant ist, also $|f| = c > 0$, $f(P) = ce^{i\theta(P)}$. Und $U_t f(P) = f(S_t P) = e^{i\lambda t} f(P)$ bedeutet $\theta(S_t P) = \theta(P) + \lambda t \dots \bmod 2\pi$. $\theta(P)$ ist offenbar überhaupt nur $\bmod 2\pi$ definiert, und (wenn man es z. B. auf Werte ≥ 0 , $< 2\pi$ reduziert) meßbar. Hat umgekehrt ein $\theta(P)$ diese

⁵⁷ Hilbert, Gött. Nachr. 1906, S. 157; Hellinger, Dissertation, Göttingen, 1907.

⁵⁸ Vgl. die Theorie der Projektionsoperatoren in E, S. 74–78.

⁵⁹ Es ist offenbar allgemein $U_t F(f(P)) = F(U_t f(P))$. Mit $f(P)$ gehört auch $|f(P)|$ zu $\mathfrak{E}$.

drei Eigenschaften, so gehört jedes $f(P) = ce^{i\theta(P)}$ zu $\mathfrak{H}$, und erfüllt $U_t f(P) = f(S_t P) = e^{i\lambda t} f(P)$, d. h. es ist Eigenfunktion vom Eigenwerte λ . Es gilt also:

Satz 4. (Koopman, a. a. O. Anm. ²) *Sei Ω eigentlich-ergodisch. $f(P)$ ist dann und nur dann Eigenfunktion vom Eigenwert λ , wenn es die Form $ce^{i\theta(P)}$ hat. Hier ist $\theta(P)$ eine mod 2π definierte, meßbare Funktion, für die $\theta(S_t P) = \theta(P) + \lambda t$ gilt.*

Zwei Eigenfunktionen f gelten als nicht wesentlich verschieden, wenn sie sich nur um einen konstanten Faktor unterscheiden; also zwei θ , wenn sie sich nur um einen konstanten Addenden unterscheiden. Für $\lambda = 0$ muß, wegen des ergodischen Charakters, $f = \text{Konstanz}$ sein, also auch $\theta = \text{Konstanz}$, d. h. im wesentlichen $\theta = 0$.

Wir wollen die $\theta(P)$ von Satz 4, mit Hinblick auf die entsprechende Begriffsbildung der klassischen Mechanik, Winkelvariable nennen, u. zw. zum Eigenwert λ gehörige.

Einige ihrer Eigenschaften liegen auf der Hand: Sind θ_1, θ_2 Winkelvariable vom Eigenwert λ , so ist $\theta_1 - \theta_2$ eine vom Eigenwert 0, d. h. konstant, also $\theta_1 = \theta_2 + \text{Konstanz}$, d. h.: zu einem gegebenen λ gehört keine oder (im wesentlichen) genau eine Winkelvariable. Ist θ eine Winkelvariable vom Eigenwert λ , so ist es $\bar{\theta}$ auch, also ist $\bar{\theta} - \theta = 2iJ\theta = \text{Konstanz}$, $J\theta = \text{Konstanz}$, d. h.: jede Winkelvariable ist (im wesentlichen) reell. Sind θ_1, θ_2 Winkelvariable der bzw. Eigenwerte λ_1, λ_2 , so ist $\theta_1 + \theta_2$ Winkelvariable des Eigenwertes $\lambda_1 + \lambda_2$. Die Winkelvariablen und ihre Eigenwerte bilden also einen „Additiv-Modul“, die letzteren sind aber das Punktspektrum der U_t : dieses enthält somit mit λ_1, λ_2 stets auch $\lambda_1 \pm \lambda_2$. Und nach dem vorher Gesagten hat die U_t -Schar nur einfache Punkteigenwerte.

Sind θ_1, θ_2 (wesentlich) verschiedene Winkelvariable (also mit verschiedenen Eigenwerten), so sind die Eigenfunktionen $e^{i\theta_1}, e^{i\theta_2}$ orthogonal, also

$$\int_{\Omega} e^{i\theta_1(P)} \overline{e^{i\theta_2(P)}} d\nu_P = \int_{\Omega} e^{i(\theta_1(P) - \theta_2(P))} d\nu_P = 0.$$

2. Gehen wir nun umgekehrt vor. Sei ein abzählbar unendliches Zahlensystem Z gegeben, das mit λ_1, λ_2 auch $\lambda_1 \pm \lambda_2$ enthält. Gibt es eine eigentlich-ergodische U_t -Schar mit reinem Punktspektrum (das, wie oben gezeigt wurde, einfach sein muß), so daß dieses Punktspektrum gleich Z ist?

Um eine solche U_t -Schar anzugeben, genügt es, ein vollständiges normiertes Orthogonalsystem beschränkter φ_{λ} (λ durchläuft Z) zu finden, das die Eigenschaft $\varphi_{\lambda} \cdot \varphi_{\mu} = \varphi_{\lambda+\mu}$ besitzt. Denn da jedes f als $\sum_{\lambda \in Z} a_{\lambda} \varphi_{\lambda}$ geschrieben werden kann, und dann $\|f\|^2 = \sum_{\lambda \in Z} |a_{\lambda}|^2$ ist (vollständiges normiertes

Orthogonalsystem!), können wir Operatoren U_t durch $U_t (\sum_{\lambda \in Z} a_\lambda \varphi_\lambda) = \sum_{\lambda \in Z} e^{i\lambda t} a_\lambda \varphi_\lambda$ definieren, dieselben sind unitär. $U_t U_s = U_{t+s}$ ist klar.

Ebenso $U_t f \cdot U_t g = U_t f \cdot g$ für $f = \varphi_\lambda$, $g = \varphi_\mu$, daher gilt es auch für alle f , die Linearaggregate von φ_λ 's sind, und deren Limites (die Beschränktheit von $g = \varphi_\mu$ erlaubt den Stetigkeitsschluß!), d. h. alle f . Ist f beschränkt, so überträgt es sich ebenso auf die Linearaggregate der φ_μ , und deren Limites (Stetigkeitsschluß!, vgl. vorhin). Somit gelten D_4 , D_5 , und nach Satz 3 entsteht die Schar U_t wirklich aus einer allgemeinen Strömung S_t . Ferner sieht man, daß $U_t \varphi_\lambda = e^{i\lambda t} \varphi_\lambda$ ist, also ist φ_λ Eigenfunktion vom Eigenwert λ . Da die φ_λ ein vollständiges normiertes Orthogonalsystem bilden, gibt es keine weiteren Eigenfunktionen. Somit hat die U_t -Schar ein reines Punktspektrum, das mit Z zusammenfällt und einfach ist. Da $0 = \lambda - \lambda$ bestimmt zu Z gehört, ist 0 einfacher Eigenwert. Die Eigenfunktion von 0, φ_0 , ist wegen $\varphi_0 \cdot \varphi_\mu = \varphi_\mu$ gleich 1, d. h. konstant. Somit ist S_t eigentlich-ergodisch, und wir sind am Ziele.

Daß aber das vollständige normierte Orthogonalsystem φ_λ , λ in Z , mit den geforderten Eigenschaften existiert, ist die Aussage eines schönen Satzes von A. Haar⁶⁰.

3. Seien nun zwei m -Räume Ω' , Ω'' mit bzw. l -Maßen μ'^* , μ''^* gegeben, und den bzw. allgemeinen Strömungen S'_t , S''_t . Diese Strömungen seien eigentlich-ergodisch, $\mu(\Omega') = \mu(\Omega'') = 1$. Die entsprechenden Operatorenscharen U'_t , U''_t mögen reine Punktspektren besitzen, und zwar beide dasselbe: Z .

Da beide Spektren einfach sein müssen, gehört zu jedem λ von Z genau je eine normierte Eigenfunktion dieses Eigenwertes bei U'_t bzw. U''_t : φ'_λ bzw. φ''_λ (diese sind nur bis auf konstante Faktoren vom Absolutwert 1 festgelegt, seien aber jetzt fest gewählt). Sowohl die φ'_λ als auch die φ''_λ bilden je ein vollständiges normiertes Orthogonalsystem.

φ'_λ hat einen konstanten Absolutwert c'_λ , also ist $\|\varphi'_\lambda\|^2 = c'^2_\lambda \cdot \mu(\Omega)$, und da die linke Seite, sowie $\mu(\Omega)$ gleich 1 ist, ist $c'_\lambda = 1$, d. h. $|\varphi'_\lambda| = 1$. Also ist auch $|\varphi'_{\lambda_1} \cdot \varphi'_{\lambda_2}| = 1$, $\varphi'_{\lambda_1} \cdot \varphi'_{\lambda_2}$ normiert. Da aber $\varphi'_{\lambda_1} \cdot \varphi'_{\lambda_2}$ Eigenfunktion vom Eigenwert $\lambda_1 + \lambda_2$ ist, muß $\varphi'_{\lambda_1} \cdot \varphi'_{\lambda_2} = c'_{\lambda_1, \lambda_2} \cdot \varphi'_{\lambda_1 + \lambda_2}$, $c'_{\lambda_1, \lambda_2}$ konstant, $|c'_{\lambda_1, \lambda_2}| = 1$, sein. Ebenso zeigt man $\varphi''_{\lambda_1} \cdot \varphi''_{\lambda_2} = c''_{\lambda_1, \lambda_2} \cdot \varphi''_{\lambda_1 + \lambda_2}$, $c''_{\lambda_1, \lambda_2}$ konstant, $|c''_{\lambda_1, \lambda_2}| = 1$.

Wir wollen nun die φ'_λ durch Multiplikation mit geeigneten Konstanten vom Absolutwerte 1 in ein neues vollständiges normiertes Orthogonal-

⁶⁰ Math. Zeitschr. 33, 1931, S. 132, Satz II, und zwar ist Z der dortigen kommutativen Gruppe $A, B, C, \dots$ gleichzusetzen, $+$ ist das Kompositionsgesetz der Gruppe, unsere $\varphi_\lambda(P)$, λ in Z , entsprechen den dortigen $X_\lambda(s)$, $X_B(s)$, $\dots$.

system von U_l -Eigenfunktionen $\psi'_\lambda = d'_\lambda \varphi'_\lambda$ überführen, so daß $k_1 \lambda_1 + \dots + k_n \lambda_n = 0$ ($\lambda_1, \dots, \lambda_n$ aus Z , $k_1, \dots, k_n = 0, \pm 1, \pm 2, \dots$) $(\psi'_{\lambda_1})^{k_1} \cdot \dots \cdot (\psi'_{\lambda_n})^{k_n} = 1$ nach sich zieht. Zu diesem Zweck ordnen wir die Elemente von Z in einer Folge an: $\lambda^{(1)}, \lambda^{(2)}, \dots$. Nehmen wir an, die $d'_{\lambda^{(1)}}, \dots, d'_{\lambda^{(m)}}$, d. h. die $\psi'_{\lambda^{(1)}}, \dots, \psi'_{\lambda^{(m)}}$ seien schon wunschgemäß festgelegt, d. h. so, daß aus $k_1 \lambda^{(1)} + \dots + k_m \lambda^{(m)} = 0$ ($k_1, \dots, k_m = 0, \pm 1, \pm 2, \dots$) $(\psi'_{\lambda^{(1)}})^{k_1} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m} = 1$ folgt; wir wollen nun $d'_{\lambda^{(m+1)}}$, d. h. $\psi'_{\lambda^{(m+1)}}$, so hinzudefinieren, daß auch aus $k_1 \lambda^{(1)} + \dots + k_m \lambda^{(m)} + k_{m+1} \lambda^{(m+1)} = 0$ ($k_1, \dots, k_m, k_{m+1} = 0, \pm 1, \pm 2, \dots$)

$$(\psi'_{\lambda^{(1)}})^{k_1} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m} \cdot (\psi'_{\lambda^{(m+1)}})^{k_{m+1}} = 1$$

folgt. Da die obige Forderung für $m = 0$ leer läuft, also erfüllt ist, können wir so induktiv alle $d'_{\lambda^{(1)}}, d'_{\lambda^{(2)}}, \dots$ und $\psi'_{\lambda^{(1)}}, \psi'_{\lambda^{(2)}}, \dots$, d. h. alle $d'_\lambda, \psi'_\lambda, \lambda$ von Z , unserer Bedingung entsprechend erzeugen. Betrachten wir also den Induktionsschritt von m zu $m+1$.

Betrachten wir die Gesamtheit der gültigen Gleichungen

$$k_1 \lambda^{(1)} + \dots + k_m \lambda^{(m)} + k_{m+1} \lambda^{(m+1)} = 0 \quad (k_1, \dots, k_m, k_{m+1} = 0, \pm 1, \pm 2, \dots).$$

Wenn $k_{m+1} = 0$ ist, so ist

$$k_1 \lambda^{(1)} + \dots + k_m \lambda^{(m)} = 0,$$

also nach Annahme

$$(\psi'_{\lambda^{(1)}})^{k_1} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m} = 1,$$

$$\text{d. h. } (\psi'_{\lambda^{(1)}})^{k_1} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m} (\psi'_{\lambda^{(m+1)}})^{k_{m+1}} = 1,$$

hier ist also, unabhängig von der Wahl von $d'_{\lambda^{(m+1)}}$, alles in Ordnung. Kommt also $k_{m+1} \neq 0$ gar nicht vor, so sind wir fertig. Kommt es vor, so sei Z_1 die Menge der k_{m+1} , die auftreten. Z_1 ist eine Menge ganzer Zahlen, es enthält offenbar mit k, l auch $k+l$, und es besteht nicht nur aus 0; somit ist Z_1 die Menge aller Vielfachen einer gewissen positiven ganzen Zahl $k^{(0)}$. Es gibt daher eine Gleichung

$$k_1^{(0)} \lambda^{(1)} + \dots + k_m^{(0)} \lambda^{(m)} + k^{(0)} \lambda^{(m+1)} = 0.$$

Es genügt

$$(\psi'_{\lambda^{(1)}})^{k_1^{(0)}} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m^{(0)}} (\psi'_{\lambda^{(m+1)}})^{k^{(0)}} = 1$$

zu sichern, denn für

$$k_1 \lambda^{(1)} + \dots + k_m \lambda^{(m)} + k_{m+1} \lambda^{(m+1)} = 0$$

ist erstens $k_{m+1} = h k^{(0)}$ (h ganz), also

$$(k_1 - h k_1^{(0)}) \lambda^{(1)} + \dots + (k_m - h k_m^{(0)}) \lambda^{(m)} = 0,$$

also nach Annahme

$$(\psi'_{\lambda^{(1)}})^{k_1 - h k_1^{(0)}} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m - h k_m^{(0)}} = 1,$$

und somit

$$(\psi'_{\lambda^{(1)}})^{k_1} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m} \cdot (\psi'_{\lambda^{(m+1)}})^{k_{m+1}} = 1.$$

Nun ist wegen

$$\varphi'_{\lambda_1} \cdot \varphi'_{\lambda_2} = c'_{\lambda_1, \lambda_2} \varphi'_{\lambda_1 + \lambda_2}, \quad |c'_{\lambda_1, \lambda_2}| = 1,$$

mit Beachtung von $\psi'_\lambda = d'_\lambda \varphi'_\lambda$, $|d'_\lambda| = 1$, jedenfalls

$$(\psi'_{\lambda^{(1)}})^{k_1^{(0)}} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m^{(0)}} = b'_1 \varphi'_{k_1^{(0)} \lambda^{(1)} + \dots + k_m^{(0)} \lambda^{(m)}} = b'_1 \varphi'_{-k^{(0)} \lambda^{(m+1)}},$$

$$(\varphi'_{\lambda^{(m+1)}})^{-k^{(0)}} = b'_2 \varphi'_{-k^{(0)} \lambda^{(m+1)}}, \quad |b'_1| = |b'_2| = 1,$$

also

$$(\psi'_{\lambda^{(1)}})^{k_1^{(0)}} \cdot \dots \cdot (\psi'_{\lambda^{(m)}})^{k_m^{(0)}} (\varphi'_{\lambda^{(m+1)}})^{k^{(0)}} = \frac{b'_1}{b'_2}.$$

Wir brauchen also bloß $d'_{\lambda^{(m+1)}}$, mit $(d'_{\lambda^{(m+1)}})^{k^{(0)}} = \frac{b'_2}{b'_1}$ zu wählen, und sind am Ziele ($|d'_{\lambda^{(m+1)}}|$ ist von selbst gleich 1).

Aus dem Bewiesenen folgt insbesondere

$$\psi'_{\lambda_1} \cdot \psi'_{\lambda_2} = \psi'_{\lambda_1 + \lambda_2} \quad (\lambda_3 = \lambda_1 + \lambda_2, \quad k_1 = k_2 = 1, \quad k_3 = -1).$$

Da wir an Stelle der ψ'_λ wieder φ'_λ schreiben können, dürfen wir annehmen, daß von vornherein $\varphi'_{\lambda_1} \cdot \varphi'_{\lambda_2} = \varphi'_{\lambda_1 + \lambda_2}$ war. Ebenso zeigt man, daß wir annehmen dürfen, daß von vornherein $\varphi''_{\lambda_1} \cdot \varphi''_{\lambda_2} = \varphi''_{\lambda_1 + \lambda_2}$ war.

Da die φ'_λ , λ in Z , und die φ''_λ , λ in Z , zwei vollständige normierte Orthogonalsysteme sind, wird durch

$$V \left(\sum_{\lambda \in Z} a_\lambda \varphi'_\lambda \right) = \sum_{\lambda \in Z} a_\lambda \varphi''_\lambda$$

ein unitärer Operator V definiert. Es ist $Vf \cdot Vg = V(f \cdot g)$ für $f = \varphi'_\lambda$, $g = \varphi'_\mu$, und hieraus schließt man, wie am Ende von 2., daß diese Gleichung für alle f, g gilt, wenn f beschränkt ist. D. h. V erfüllt D_ε . Auf Grund des in § III, 1. Bewiesenen muß daher eine eindeutige und maßtreue Abbildung von Ω' minus einer Menge vom Maße 0 auf Ω'' minus eine Menge vom Maße 0 existieren, die V nach $Vf(P) = f(\Sigma P)$ erzeugt.

Ferner ist

$$U_t'' \varphi''_\lambda = e^{it\lambda} \varphi''_\lambda, \quad V U_t' V^{-1} \varphi''_\lambda = V U_t' \varphi'_\lambda = V e^{it\lambda} \varphi'_\lambda = e^{it\lambda} \varphi''_\lambda,$$

$$\text{d. h. es ist } U_t'' f = V U_t' V^{-1}$$

für alle $f = \varphi'_\lambda$. Daher gilt dies auch für die Linearaggregate der φ'_λ und deren Häufungspunkte, d. h. für alle f , also ist $U_t'' = V U_t' V^{-1}$. Mit Hilfe der S_t' , S_t'' , Σ ausgedrückt, besagt dies:

$$f(S_t'' P'') = f(\Sigma S_t' \Sigma^{-1} P'') \quad (P'' \text{ in } \Omega''),$$

und da dies für alle f gilt, muß $S_t'' \cong \Sigma S_t' \Sigma^{-1}$ sein. D. h. die zwei allgemeinen Strömungen S_t' und S_t'' sind isomorph.

4. Das bisher Bewiesene zusammenfassend, können wir also aussagen:

SATZ 5. Wenn $\mu(\Omega) = 1$ und die allgemeine Strömung S_t in Ω (eigentlich-)ergodisch ist, und wenn die Schar U_t ein reines Punktspektrum hat, so ist dieses eine abzählbar-unendliche Menge Z von reellen Zahlen, die mit λ , μ auch $\lambda \pm \mu$ enthält.

Umgekehrt werden auf diese Weise alle abzählbar-unendlichen Mengen Z von reellen Zahlen, die mit λ , μ auch $\lambda \pm \mu$ enthalten, erzeugt, und zwar bestimmt Z Ω , S_t im wesentlichen (d. h. bis auf isomorphe Abbildungen, vgl. D_3) eindeutig.

Jede Zahlenmenge Z der beschriebenen Art, d. h. jeder „Zahlenmodul“, bestimmt somit im wesentlichen eindeutig eine eigentlich-ergodische Strömung. Damit ist die Existenz unendlich vieler Typen ergodischer Strömungen bewiesen.

Indessen ist es lehrreich, für die einfachsten „Zahlenmoduln“ Z Beispiele dazugehöriger Strömungen (die ja im wesentlichen die einzigen sind) effektiv anzugeben.

Wenn z. B. Z eine endliche Basis hat — d. h. wenn ein Zahlensystem $\omega_1, \dots, \omega_n$ existiert, so daß erstens keine Relation $k_1 \omega_1 + \dots + k_n \omega_n = 0$ ($k_1, \dots, k_n = 0, \pm 1, \pm 2, \dots$), außer mit $k_1 = \dots = k_n = 0$, besteht, und zweitens Z die Menge aller $k_1 \omega_1 + \dots + k_n \omega_n$ ($k_1, \dots, k_n = 0, \pm 1, \pm 2, \dots$) ist — so werden wir zu einer wohlbekannten, von Weyl entdeckten Strömung geführt⁶¹.

Sei nämlich Ω der n -dimensionale Würfel, $P = [u_1, \dots, u_n]$, $0 \leq u_1 < 1, \dots, 0 \leq u_n < 1$, u^* das gewöhnliche Lebesguesche äußere Maß. Es sei

$$S_t[u_1, \dots, u_n] = \left[\left\{ u_1 + \frac{\omega_1}{2\pi} t \right\}, \dots, \left\{ u_n + \frac{\omega_n}{2\pi} t \right\} \right]$$

($\{u\}$ ist der Rest von u mod 1: $0 \leq \{u\} < 1$, $u - \{u\}$ ganz). Die Funktionen $e^{2\pi i(k_1 u_1 + \dots + k_n u_n)}$ ($k_1, \dots, k_n = 0, \pm 1, \pm 2, \dots$) bilden bekanntlich ein vollständiges normiertes Orthogonalsystem; und wegen

$$S_t e^{2\pi i(k_1 u_1 + \dots + k_n u_n)} = e^{i(k_1 \omega_1 + \dots + k_n \omega_n) t} e^{2\pi i(k_1 u_1 + \dots + k_n u_n)}$$

⁶¹ Weyl,

sind es Eigenfunktionen mit den bzw. Eigenwerten $k_1 \omega_1 + \dots + k_n \omega_n$. Also gibt es keine weiteren Eigenfunktionen. Somit liegt ein reines Punktspektrum vor, das gleich Z ist; und 0 ist einfacher Eigenwert mit der Eigenfunktion 1, d. h. die Strömung ist eigentlich ergodisch. Dieses Ω und S_t ist also das gesuchte Beispiel.

Es sei noch erwähnt, daß wir hier für $n = 1$ die Menge der $k\omega$ ($k = 0, \pm 1, \pm 2, \dots$) als Z erhalten, und S_t die einfach periodische Bewegung mit der Periode $\frac{2\pi}{\omega}$ wird. Dieses Z ist der einzige nicht überall dichte „Zahlmodul“⁶², in allen anderen Fällen mit reinem Punktspektrum ist also dieses überall dicht.

Ein etwas komplizierterer Fall ist dieser: Z ist die Menge aller dyadisch rationalen Zahlen $\left(\frac{n}{2^p}, n = 0, \pm 1, \pm 2, \dots, p = 1, 2, \dots\right)$, denn hier existiert keine Basis. Ω sei das Quadrat $P = [u, v], 0 \leq u < 1, 0 < v < 1$, μ^* sei das gewöhnliche Lebesguesche äußere Maß. S_t werde so definiert: $\varphi(v)$ ist die folgende eindeutige und maßtreue Abbildung von $0 < v < 1$ auf sich selbst: $\frac{1}{2} < v < 1$ wird um $-\frac{1}{2}$ nach $\frac{1}{4} < \varphi(v) < \frac{1}{2}$ verschoben, $\frac{1}{4} < v < \frac{1}{2}$ um $+\frac{1}{4}$ nach $\frac{1}{2} < \varphi(v) < \frac{3}{4}$, $\frac{1}{8} < v < \frac{1}{4}$ um $+\frac{5}{8}$ nach $\frac{3}{4} < v < \frac{7}{8}$, ..., die übrigbleibenden Punkte $v = \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \dots$ werden in dieser Reihenfolge den übrigbleibenden $\varphi(v) = \frac{1}{2}, \frac{3}{4}, \frac{7}{8}, \dots$ zugeordnet⁶³. $S_t[u, v]$ entsteht, indem man u um $\frac{t}{2\pi}$ wachsen läßt (d. h. für $t > 0$ nach rechts, für $t < 0$ nach links verschiebt), aber so oft man hierbei den rechten oder linken Rand von Ω (d. i. $u = 1$ bzw. $u = 0$) passiert, u auf den linken bzw. rechten Rand zurückversetzt und v durch $\varphi(v)$ bzw. $\varphi^{-1}(v)$ (Inverse!) ersetzt⁶⁴. Das Spektrum der U_t -Schar bestimmt man am bequemsten so: Man versuche eine in jedem der 2^p Intervalle $\frac{m-1}{2^p} < v < \frac{m}{2^p}$

⁶² Z hat Elemente $\lambda \neq 0$, also auch Elemente $> 0: |\lambda| = \pm \lambda$. Gibt es unter seinen Elementen > 0 ein kleinstes, λ_0 , so schließen wir so: für jedes λ von Z gibt es ein $k = 0, \pm 1, \pm 2, \dots$ mit $\lambda - k\lambda_0 \geq 0, < \lambda_0$, also kann nicht $\lambda - k\lambda_0 > 0$ sein, also ist $\lambda - k\lambda_0 = 0, \lambda = k\lambda_0$. D. h. ist die im Text genannte Menge. Andernfalls wird die untere Grenze der $\lambda > 0$ von Z , λ_0 , nicht angenommen. Zu $\epsilon > 0$ gibt es ein λ_1 in Z mit $\lambda_0 < \lambda_1 < \lambda_0 + \epsilon$, und ein λ_2 in Z mit $\lambda_0 < \lambda_2 < \lambda_1$. Also ist das $\lambda = \lambda_1 - \lambda_2$ von Z $> 0, < \epsilon$. Bei beliebigem μ können wir daher $\mu < k\lambda < \mu + \epsilon$ erreichen, da $\mu, \epsilon > 0$ beliebig sind, ist Z überall dicht.

⁶³ In Formeln: $\varphi(v) = v + 1 - \frac{3}{2^n}$ für $\frac{1}{2^n} < v < \frac{1}{2^{n-1}}$, und $\varphi(v) = 1 - v$ für $v = \frac{1}{2^n}$ ($n = 1, 2, \dots$).

⁶⁴ In Formeln: $\varphi^n(v)$ sei die $|n|$ -fache Iterierte von $\varphi(v)$ bzw. $\varphi^{-1}(v)$ (für $n > 0$ bzw. < 0), dann ist $S_t[u, v] = \left[\left\{ u + \frac{t}{2\pi} \right\}, \varphi^{u + \frac{t}{2\pi} - \left\{ u + \frac{t}{2\pi} \right\}}(v) \right] \left\{ u + \frac{t}{2\pi} - \left\{ u + \frac{t}{2\pi} \right\} \right\}$ ist ganz.

($m = 1, 2, \dots, 2^p$) konstante Funktion $\sigma_{n,p}(v)$ so zu bestimmen, daß $\varphi_{n,p}(u, v) = \sigma_{n,p}(v) e^{\frac{2\pi i}{2^p} n u}$ eine normierte Eigenfunktion vom Eigenwert $\frac{n}{2^p}$ wird. Es zeigt sich, daß dies für jedes $n = 0, \pm 1, \pm 2, \dots$, $p = 1, 2, \dots$, auf genau eine Weise gelingt (bis auf einen unwesentlichen konstanten Faktor). Da $\varphi_{n,p}$ den an $\varphi_{2n,p+1}$ gestellten Anforderungen genügt, ist $\varphi_{n,p} = \varphi_{2n,p+1}$; also $\varphi_{n,p}$ nur vom Werte von $\lambda = \frac{n}{2^p}$ abhängig, also auch mit φ_λ bezeichnbar. Da sie Eigenfunktionen sind, bilden die φ_λ ein normiertes Orthogonalsystem, es zeigt sich, daß es ein vollständiges ist. Also gibt es keine weiteren Eigenfunktionen. Somit liegt ein reines Punktspektrum vor, das gleich Z ist; und 0 ist einfacher Eigenwert, natürlich mit der Eigenfunktion 1, d. h. die Strömung ist eigentlich ergodisch. Ω, S_t sind also das gesuchte Beispiel. Die einfache Durchführung der zwei weiter oben ausgelassenen Details übergehen wir.

Nach ähnlichen Schemata, wie diese zwei, ließen sich noch viele andere „Zahlmoduln“ Z behandeln, worauf wir nicht mehr eingehen.

V. Bemerkungen über U_t -Scharen mit Streckenspektrum.

1. Die U_t -Scharen mit reinem Streckenspektrum (d. h. vom eventuellen einfachen Eigenwert 0 abgesehen, ohne Punktspektrum) sind, wie schon in Einleitung 4 betont wurde, aus zwei Gründen von besonderer Wichtigkeit. Erstens die Verallgemeinerung des Ergodensatzes (e) in Einleitung 2, vgl. a. a. O. Anm.³) für sie und nur für sie, d. h. sie sind, wie a. a. O. des näheren erläutert wurde, die Strömungen ohne „à la longue“ gültige Eigenschaften. Zweitens ist zu vermuten, daß die U_t -Scharen im allgemeinen ein reines Streckenspektrum haben, d. h. daß das Auftreten eines Punktspektrums eine seltene Ausnahmeerscheinung ist. Dies letztere werden wir im § VI plausibel zu machen suchen. Auf Grund dieser Umstände wäre es erwünscht, sie mindestens so gut zu übersehen, wie dies in § V für die Punktspektren gelang.

Zunächst wäre zu erwarten, daß das Streckenspektrum, wie das Punktspektrum, stets einfach ist (ergodische Strömung vorausgesetzt!). Weiter, daß ein Isomorphiesatz gilt, wie er in Einleitung 3. formuliert und in der zweiten Hälfte von Satz 5 für das reine Punktspektrum bewiesen wurde. D. h. daß zwei allgemeine Strömungen Ω, S_t mit unitär ineinander transformierbaren U_t einander im wesentlichen (d. h. bis auf eine eindeutige und maßtreue Abbildung) gleich sind. Dann wüßten wir, was alle für uns bedeutungsvollen Eigenschaften einer Strömung (d. s. die Isomorphie-Invarianten, vgl. Einleitung 3.) festlegt: so wie dies im Falle des Punkt-

spektrums nach Satz 5 das Punktspektrum Z war, wäre es im Falle des Streckenspektrums dieses und sein Hellingerscher Typus (vg. Anm. ²⁵). Hinzu tritt noch, daß wir, wenigstens für differentiierbare Strömungen, zeigen werden (und es ist für alle zu vermuten): wenn ein Streckenspektrum überhaupt existiert, so umfaßt es die ganze Zahlengerade $-\infty < \lambda < +\infty$. Also käme es nur auf den Hellingerschen Typus an.

Schließlich wissen wir aus Satz 5, daß das Punktspektrum Z nicht jede beliebige (abzählbar-unendliche) Zahlenmenge sein kann: vielmehr ist die Beschränkung da, daß Z ein „Zahlmodul“ sein muß, aber keine weitere. Es wäre also wichtig, zu erfahren, welchen Beschränkungen der Hellingersche Typus unterliegt. Wenn er z. B. immer der einfachste Typus („ λ “) wäre, so würde dies eine bemerkenswerte Verschärfung von e) (Einleitung 2., vgl. die zweite in Anm. ³ genannte Abhandlung, ferner Anm. ²⁵) nach sich ziehen. Ferner hätte es die überraschende Folge, daß alle Strömungen mit reinem Streckenspektrum (d. h. gerade der vermutlich allgemeine Fall) einander im wesentlichen gleich (d. h. isomorph) wären. Es ist aber wohl zu gewagt, hierüber Vermutungen anzustellen.

Von alledem können wir leider einstweilen nichts beweisen. Bloß die Ausdehnung des Streckenspektrums bei differentiierbaren Strömungen soll im Folgenden bestimmt werden.

2. Ω sei ein m -Raum, der „im kleinen“ auf n -dimensionale Koordinaten bezogen werden kann (d. h. in einer geeigneten Umgebung eines jeden P_0 $P = \{u_1, \dots, u_n\}$), und S_t eine ergodische und differentiierbare Strömung in Ω .

Wenn P ein Punkt von Ω ist, so sei $S(P)$ die Bahnkurve $S_t P$, $-\infty < t < +\infty$. Wenn $S_t P$ Doppelpunkte hat, d. h. $S_{t_1} P = S_{t_2} P$, $t_1 < t_2$, vorkommt, so ergibt das Anwenden von S_{t-t_1} : $S_t P = S_{t+(t_2-t_1)} P$, d. h. es gibt eine positive Periode $t_2 - t_1$, die kleinste solche Periode sei τ . Tritt dies für jedes P ein, so ist $\tau = \tau(P) > 0$, wie man leicht einsieht, eine meßbare Funktion von P . Es gibt also ein P_0 , wo sie approximativ stetig ist (dies ist sogar für alle P_0 bis auf eine Menge vom Maße 0 der Fall ⁶⁵). Sei K eine Umgebung von P_0 , K_1 die Menge der $S_t P$, P in K , $0 \leq t \leq \tau(P_0) + 1$. Das Maß dieser Menge ist offenbar endlich, und da sie bei sich auf P_0 zusammenziehendem K auf den Bahnkurvenbogen $S_t P_0$, $0 \leq t \leq \tau(P_0) + 1$, d. h. $S(P_0)$, zusammenzieht (weil $S_t P$ stetig ist!), strebt das Maß gegen das Maß von $S(P_0)$. Ist dieses $< \mu(\Omega)$, so können wir K mit $\mu(K_1) < \mu(\Omega)$ wählen. Da $\tau(P)$ bei $P = P_0$ approximativ stetig ist, ist in einer Menge $\mathcal{A} \subseteq K$, $\mu(\mathcal{A}) > 0$, $\tau(P) \leq \tau(P_0) + 1$ (es ließe sich noch mehr behaupten). Sei $\mathcal{A}_1$ die Menge der $S_t P$, P in $\mathcal{A}$, $0 \leq t \leq \tau(P_0) + 1$, wegen $\tau(P) \leq \tau(P_0) + 1$ ist $\mathcal{A}_1$ die Summe der $S(P)$, P in $\mathcal{A}$, also S_t -invariant. Dabei ist

⁶⁵ Vgl. z. B. Sierpinski, Fund. Math. 3, 1922, S. 320.

$$\mu(\mathcal{A}) \geq \mu(\mathcal{A}_1) > 0, \quad \leq \mu(K_1) < \mu(\Omega).$$

Somit wäre S_t nicht ergodisch.

Also muß $\mu(S(P_0)) = \mu(\Omega)$ sein. Da wir eine Menge vom Maße 0 fortlassen können, können wir $S(P_0) = \Omega$ annehmen. $S(P_0)$ ist eineindeutig auf die $s \geq 0, \leq \tau (= \tau(P_0))$ bezogen (durch $s \rightarrow S_s P_0$). Wir können Ω somit durch die Menge $0 \leq t < \tau$ ersetzen unter Übertragung des Maßes. Da

$$S_t S_s = S_{t+s} = S_{t+s-\epsilon\tau} \quad (e = 0, \pm 1, \pm 2, \dots)$$

ist, ist nunmehr S_t durch $S_t s = t + s \dots \bmod \tau$ definiert. Daraus, daß das Maß dieser Operation gegenüber invariant sein muß, folgert man unschwer, daß es das Lebesguesche ist, eventl. mit einer Konstanten $a > 0$ multipliziert. Und nunmehr sieht man, daß die Funktionen $\frac{1}{\sqrt{a\tau}} e^{\frac{2\pi i}{\tau} ks}$, $k = 0, \pm 1, \pm 2, \dots$, Eigenfunktionen der U_t sind, mit den bzw. Eigenwerten $\frac{2\pi}{\tau} k$; da sie ein vollständiges normiertes Orthogonalsystem bilden, gibt es keine anderen Eigenfunktionen. Es liegt also ein reines Punktspektrum vor, $Z: \frac{2\pi}{\tau} k$, $k = 0, \pm 1, \pm 2, \dots$. Das ist gerade der in § IV, 4. diskutierte allereinfachste Fall: einfach periodische Bewegung. Insbesondere ist $n = 1$.

Wenn dieser nicht vorliegt, so muß es demnach ein P_0 mit einfach durchlaufener Bahnkurve $S(P_0)$ geben: alle $S_t P_0$ voneinander verschieden. Wir wählen nun ein t_0 , das wir vorläufig festhalten. Ein quer zu $S(P_0)$ liegendes, $n-1$ -dimensionales Flächenelement durch P_0 in Ω , F , können wir so wählen, daß alle $S_t F$, $0 \leq t \leq t_0$ zueinander fremd sind (Differenzierbarkeit!, F muß nur klein genug sein). Sei $F(t_1, t_2)$ ($t_1 < t_2$) die Menge der $S_t P$, P in F , $t_1 \leq t \leq t_2$. Da $F(t_1, t_2)$ ein n -dimensionaler Bereich ist, ist $\mu(F(t_1, t_2))$ stets > 0 . Aus der Maßtreue von S_t folgt

$$\mu(F(t_1, t_2)) = \mu(S_t F(t_1, t_2)) = \mu(F(t_1 + t, t_2 + t)),$$

d. h.

$$\mu(F(t_1, t_2)) = f(t_2 - t_1).$$

Liegen t_1, t_2, t_3 in $0 \leq t \leq t_0$, so ist

$$\mu(F(t_1, t_2)) + \mu(F(t_2, t_3)) = \mu(F(t_1, t_3)),$$

d. h. es ist

$$f(s') + f(s'') = f(s' + s''),$$

wenn $s', s'' > 0$, $s' + s'' \leq t_0$ ist. Da ferner $f(s)$ stets > 0 ist (es ist ja nur für $s > 0$, $\leq t_0$ definiert), erkennt man, daß $f(s) = cs$, $c = \text{Konstanz} > 0$, sein muß. Also $\mu(F(t_1, t_2)) = c(t_1 - t_2)$.

Wir definieren nun: $\varphi(P) \begin{cases} = e^{i\lambda s}, & \text{wenn } P \text{ in } S_s F \text{ liegt, } 0 \leq s \leq t_0 \\ = 0, & \text{sonst (d. h. } P \text{ nicht in } F(0, t_0)) \end{cases}$.
Es ist

$$\|\varphi\|^2 = \int_{F(0, t_0)} ds = \mu(F(0, t_0)) = ct_0,$$

$$\|U_t \varphi_0 - e^{i\lambda t} \varphi_0\| = \int_{F(0, t) + F(t_0, t_0+t)} ds = \mu(F(0, t)) + \mu(F(t_0, t_0+t)) = 2ct,$$

für $0 \leq t \leq t_0$. Nun ist $(E(\lambda) - E(-\lambda))\varphi \rightarrow \varphi$ für $\lambda \rightarrow \infty$, also für ein geeignetes λ

$$\|(E(\lambda) - E(-\lambda))\varphi\| \geq \frac{1}{2} ct_0.$$

Andererseits ist $E(\lambda) - E(-\lambda)$ ein Projektionsoperator, und mit U_t vertauschbar, also

$$\begin{aligned} & \|U_t (E(\lambda) - E(-\lambda))\varphi - e^{i\lambda t} (E(\lambda) - E(-\lambda))\varphi\| \\ &= \|(E(\lambda) - E(-\lambda))(U_t \varphi - e^{i\lambda t} \varphi)\| \leq \|U_t \varphi - e^{i\lambda t} \varphi\| = 2ct. \end{aligned}$$

Für $\psi = \frac{(E(\lambda) - E(-\lambda))\varphi}{\|(E(\lambda) - E(-\lambda))\varphi\|}$ haben wir also:

$$\|\psi\| = 1, \quad \|U_t \psi - e^{i\lambda t} \psi\| \leq \frac{4t}{t_0}.$$

Ferner ist A auf $(E(\lambda) - E(-\lambda))\varphi$ anwendbar⁶⁶, also auch auf ψ . Unsere Ungleichheit kann auch so geschrieben werden:

$$\left\| \frac{U_t - 1}{it} \psi - \frac{e^{i\lambda t} - 1}{it} \psi \right\| \leq \frac{4}{t_0},$$

und hieraus folgt durch $t \rightarrow 0$:

$$\|A\psi - \lambda\psi\| \leq \frac{4}{t_0}.$$

t_0 war beliebig, wir können es also, wenn ein beliebiges $\epsilon > 0$ gegeben ist, mit $\frac{4}{t_0} < \epsilon$ wählen. Dann wird $\|A\psi - \lambda\psi\| < \epsilon \|\psi\|$. Da für jedes $\epsilon > 0$ ein solches λ gefunden werden kann, gehört λ zum Gesamtspektrum. D. h.: das Gesamtspektrum ist die ganze Zahlengerade $-\infty < \lambda < +\infty$.

3. Wenn kein Punktspektrum da ist, oder das Punktspektrum Z die Menge der $k\omega$, $k = 0, \pm 1, \pm 2, \dots$, für ein festes $\omega > 0$ ist, so enthält

⁶⁶ Vgl. z. B. A, S. 415, Anm. ⁶⁸

das Streckenspektrum alle Punkte des Gesamtspektrums, bis auf höchstens die $k\omega$. Wenn also nicht der in 2. erwähnte Spezialfall (reines Punktspektrum, $Z = \text{Menge der } k\omega, k = 0, \pm 1, \pm 2, \dots$) vorliegt, so enthält das Gesamtspektrum alle λ , so daß das Streckenspektrum überall dicht ist. Da es eine abgeschlossene Menge ist, enthält es alle λ .

Andernfalls ist das Punktspektrum Z überall dicht (vgl. Anm. ⁶²). Gehöre λ zu Z , φ_λ sei eine Eigenfunktion vom Eigenwert λ . Nach Satz 4 ist $|\varphi_\lambda| = \text{Konstanz} > 0$, wir können also $|\varphi_\lambda| = 1$ annehmen. Der Operator $Vf = \varphi_\lambda \cdot f$ ist also unitär. Aus $U_t \varphi_\lambda = e^{it\lambda} \varphi_\lambda$ folgt

$$V^{-1} U V = e^{it\lambda} U.$$

Nun hat $V^{-1} U V$ dasselbe Spektrum wie U , $e^{it\lambda} U$ dagegen offenbar ein um λ nach rechts verschobenes. Insbesondere das Streckenspektrum von U geht also durch eine Verschiebung nach rechts um λ in sich selbst über. Wenn also μ dazu gehört, gehören alle $\mu + \lambda$, λ in Z , auch dazu, es ist also überall dicht. Da es eine abgeschlossene Menge ist, enthält es dann alle μ . D. h.: es ist leer, oder die ganze Zahlengerade $-\infty < \lambda < +\infty$.

Zusammenfassend können wir also sagen:

Satz 5. *Für eine ergodische differentiierbare Strömung S_t ist das Gesamtspektrum immer die ganze Zahlengerade $-\infty < \lambda < +\infty$, mit einer Ausnahme: bei der einfach periodischen Strömung mit der Periode $\frac{2\pi}{\omega}$, $\omega > 0$, und dem reinen Punktspektrum $Z: k\omega, k = 0, \pm 1, \pm 2, \dots$.*

Das Streckenspektrum fehlt, oder es ist die ganze Zahlengerade $-\infty < \lambda < +\infty$.

Dieser Satz gilt wahrscheinlich auch unabhängig von der Annahme der Differentiierbarkeit für die Strömung.

VI. Beispiele für ein reines Streckenspektrum.

1. Aus § IV, 1. wissen wir, daß jede nullergodische Strömung ein reines Streckenspektrum hat. Solche Strömungen sind leicht anzugeben: z. B. Ω die Menge aller reellen μ , μ^* das gewöhnliche Lebesguesche äußere Maß, $S_t \mu = \mu + t$. Denn wenn $\mathcal{A}$ meßbar und von endlichem Maße ist, so können wir $t_0 > 0$ so wählen, daß der in $-t_0 < u < t_0$ liegende Teil von $\mathcal{A}$ ein Maß $> \frac{1}{2} \mu(\mathcal{A})$ hat. Der in $-3t_0 < u < -t_0$ liegende hat daher ein Maß $< \frac{1}{2} \mu(\mathcal{A})$, also auch der in $-t_0 < u < t_0$ liegende Teil von $S_{2t_0} \mathcal{A}$, so daß sich $\mathcal{A}, S_{2t_0} \mathcal{A}$ nicht nur um eine Menge vom Maße 0 unterscheiden können. Jede $\mathcal{A}$ -Menge muß daher das Maß 0 haben, d. h. S_t ist nullergodisch.

Uns interessieren aber in erster Linie die eigentlich-ergodischen Strömungen, und daher wollen wir für solche Beispiele mit reinem Streckenspektrum angeben.

Die zu diskutierenden Beispiele werden alle von dieser Art sein: $f(u)$ sei eine in $0 \leq u < 1$ definierte Funktion, die stets $\geq A, \leq B$ bleibt, für zwei feste positive A, B ; γ eine Zahl, $0 \leq \gamma < 1$. Ω sei der ebene Bereich $P = [u, v]$, $0 \leq u < 1$, $0 \leq v < f(u)$, μ^* das gewöhnliche Lebesguesche äußere Maße. $S_t[u, v]$ entsteht, indem man v um t wachsen läßt (für $t > 0$, für $t < 0$ ist das eine Abnahme), aber jedesmal, wenn man die obere oder die untere Grenze von Ω ($v = f(u)$ bzw. $v = 0$) passiert, v auf die untere bzw. obere Grenze von Ω zurückversetzt, und u durch $\{u + \gamma\}$ bzw. $\{u - \gamma\}$ ersetzt (d. h. durch diejenige der Zahlen $u + \gamma$, $u + \gamma - 1$ bzw. $u - \gamma$, $u - \gamma + 1$, die $\geq 0, < 1$ ist). S_t ist offenbar maßtreu.

Wir werden nun untersuchen, für welche $f(u), \gamma$ diese Strömung (eigentlich-) ergodisch ist, und ein reines Streckenspektrum hat.

2. Gehöre λ zum Punktspektrum der zu den S_t gehörigen Operatorenschar U_t , $\varphi \neq 0$ sei eine dazugehörige Eigenfunktion. D. h. es soll für jedes feste t , bis auf eine $[u, v]$ -Menge vom Maße 0, $\varphi(S_t[u, v]) = e^{i\lambda t} \varphi(u, v)$ sein. Wir werden nun mit Hilfe des Umstandes, daß S_t nicht nur D'_2 , sondern auch D_2 erfüllt, die obigen Ausnahmемengen eliminieren.

Da $\varphi(u, v)$ in $[u, v]$ meßbar ist, und das $S_t[u, v]$ -Urbild jeder meßbaren Menge eine meßbare $[u, v, t]$ -Menge ist (vgl. § II, 6), ist $\varphi(S_t[u, v])$ in $[u, v, t]$ meßbar. Die $[u, v, t]$ -Menge $\bar{N}$, die durch $\varphi(S_t[u, v]) \neq e^{i\lambda t} \varphi(u, v)$ definiert ist, ist also meßbar. Für jedes feste t ist die Menge $\bar{N}_{(t)}$ der $[u, v]$ mit $[u, v, t]$ in $\bar{N}$ eine Menge vom Maße 0. Nach dem Satz von Fubini ist also auch $\bar{N}$ vom Maße 0. Nach demselben Satz (in umgekehrter Richtung angewandt), ist für alle $[u, v]$, abgesehen von einer $[u, v]$ -Menge N vom Maße 0, die Menge der t mit $[u, v, t]$ in $\bar{N}$ vom Maße 0. D. h.: wenn $[u, v]$ nicht zu N gehört, gilt $\varphi(S_t[u, v]) = e^{i\lambda t} \varphi(u, v)$ für alle t , bis auf eine Menge vom Maße 0.

Nun sollen $[u', v']$, $[u'', v'']$ nicht zu N gehören, ferner sei

$$S_{t'}[u', v'] = S_{t''}[u'', v''].$$

Es gilt

$$\varphi(S_{s'}[u', v']) = e^{i\lambda s'} \varphi(u', v'), \quad \varphi(S_{s''}[u'', v'']) = e^{i\lambda s''} \varphi(u'', v''),$$

ferner

$$\varphi(S_t S_{t'}[u', v']) = e^{i\lambda t} \varphi(S_{t'}[u', v']), \quad \varphi(S_t S_{t''}[u'', v'']) = e^{i\lambda t} \varphi(S_{t''}[u'', v'']),$$

bis auf je eine Menge vom Maße 0, bzw. in s', s'', t, t' . Binden wir s', s'', t durch $s' = t + t'$, $s'' = t + t''$ aneinander, so gelten alle vier Gleichungen bis auf eine t -Menge vom Maße 0. Wählen wir also t so,

— Oder, für ein festes positives c , das c -fache davon. Dies würde am folgenden nichts ändern, und wir könnten es benützen, um $\mu(\Omega) = c \int_0^1 f(u) du$ gleich 1 zu machen.

daß sie gelten, und beachten $S_t S_{t'} = S_{s'}$, $S_t S_{t''} = S_{s''}$ sowie $S_{t'} [u', v'] = S_{t''} [u'', v'']$, so wird

$$e^{i\lambda s'} \varphi(u', v') = e^{i\lambda s''} \varphi(u'', v''), \quad e^{i\lambda t'} \varphi(u', v') = e^{i\lambda t''} \varphi(u'', v'').$$

Wenn wir also für alle $[\bar{u}, \bar{v}] = S_t [u, v]$, $[u, v]$ nicht in N , ein neues $\bar{\varphi}(\bar{u}, \bar{v})$ durch $e^{i\lambda t} \varphi(u, v)$ definieren, so ist diese Definition nicht selbstwidersprechend. Für die übrigen $[\bar{u}, \bar{v}]$ sei etwa $\bar{\varphi}(\bar{u}, \bar{v}) = 0$. Gehört $[\bar{u}, \bar{v}]$ nicht zu N , so können wir $u = \bar{u}$, $v = \bar{v}$, $t = 0$ setzen, und es wird $\bar{\varphi}(\bar{u}, \bar{v}) = \varphi(\bar{u}, \bar{v})$. Also tritt $\bar{\varphi}(u, v) \neq \varphi(u, v)$ nur in N , also in eine Menge vom Maße 0 ein, d. h. $\bar{\varphi}$ gehört, wie φ , zu $\mathfrak{H}$, ja es ist als $\mathfrak{H}$ -Element gleich φ . Ferner gilt nach Konstruktion ausnahmslos

$$\bar{\varphi}(S_t [u, v]) = e^{i\lambda t} \bar{\varphi}(u, v).$$

Anstatt φ durch $\bar{\varphi}$ zu ersetzen, können wir auch annehmen, daß φ von vornherein so war.

3. Wenn $[u, v']$, $[u, v'']$ beide in $\mathfrak{L}$ liegen, so ist für $t = v'' - v'$ $S_t [u, v'] = [u, v'']$, also

$$\varphi(u, v'') = e^{i\lambda(v'' - v')} \varphi(u, v'), \quad e^{-i\lambda v'} \varphi(u, v') = e^{-i\lambda v''} \varphi(u, v'').$$

$e^{-i\lambda v} \varphi(u, v)$ ist also von v unabhängig, es ist somit eine meßbare u -Funktion: $\psi(u)$. Sei nun $0 < t < 2A$, $v = f(u) - \frac{t}{2}$, dann ist

$$S_t \left[u, f(u) - \frac{t}{2} \right] = \left[\{u + \gamma\}, \frac{t}{2} \right],$$

also

$$e^{i\lambda \frac{t}{2}} \psi(\{u + \gamma\}) = e^{i\lambda t} \cdot e^{i\lambda \left(f(u) - \frac{t}{2}\right)} \varphi(u), \quad \psi(\{u + \gamma\}) = e^{i\lambda f(u)} \varphi(u).$$

Umgekehrt erkennt man: wenn $\psi(u)$ meßbar ist und

$$\psi(\{u + \gamma\}) = e^{i\lambda f(u)} \psi(u)$$

gilt, so ist

$$\varphi(u, v) = e^{i\lambda v} \psi(u)$$

eine Lösung von

$$\varphi(S_t [u, v]) = e^{i\lambda t} \varphi(u, v).$$

Da ferner

$$\begin{aligned} \int \int_{\mathfrak{L}} |\varphi(u, v)|^2 du dv &= \int_0^1 \int_0^{f(u)} |\psi(u)|^2 du dv \\ &= \int_0^1 |\psi(u)|^2 f(u) du \left\{ \begin{array}{l} \geq A \cdot \int_0^1 |\psi(u)|^2 du \\ \leq B \cdot \int_0^1 |\psi(u)|^2 du \end{array} \right. \end{aligned}$$

ist, ist die Erfüllbarkeit der obigen Bedingungen für $\psi(u)$, zusammen mit der Endlichkeit von $\int_0^1 |\psi(u)|^2 du$, für den Eigenwert-Charakter von λ charakteristisch.

Für $\lambda = 0$ besagt dies $\psi(\{u + \gamma\}) = \psi(u)$. Nun bilden die $e^{2\pi i k u}$, $k = 0, \pm 1, \pm 2, \dots$, ein vollständiges normiertes Orthogonalsystem, und sind Eigenfunktionen des unitären Operators $I\chi(u) = \chi(\{u + \gamma\})$, und zwar mit den bzw. Eigenwerten $e^{2\pi i k \gamma}$. Also gibt es keine weiteren Eigenfunktionen. Ist γ rational, so sind viele $e^{2\pi i k \gamma} = 1$, also viele $\psi(u) = e^{2\pi i k u}$ möglich, d. h. unsere Strömung nicht ergodisch. Ist es dagegen irrational, was im folgenden vorausgesetzt werde, so folgt aus $e^{2\pi i k \gamma} = 1$ $k = 0$, also muß $\psi(u)$ ein konstantes Vielfaches der einzigen Eigenfunktion, 1, sein — also $\psi(u) = c$, $\varphi(u, v) = c$, d. h. unsere Strömung ist ergodisch.

Sei nun $\lambda \neq 0$. Es ist $|\psi(\{u + \gamma\})| = |\psi(u)|$, also die obigen Schlüsse auf $|\psi(u)|$ anwendbar: es ist $|\psi(u)| = c$, wegen $\varphi \neq 0$, $\psi \neq 0$ ist $c > 0$.

$\int_0^1 |\chi(\{u + \varepsilon\}) - \chi(u)|^2 du \rightarrow 0$ für $\varepsilon \rightarrow 0$ ist eine wohlbekannte Tatsache⁶⁹. — Nach dem Satze von Kronecker⁷⁰ können wir, da γ irrational ist, eine Folge ganzer Zahlen $k_1 < k_2 < \dots$ mit $\{k_n \gamma\} \rightarrow 0$ für $n \rightarrow \infty$ finden. Also ist dann

$$\int_0^1 |\psi(\{u + k_n \gamma\}) - \psi(u)|^2 du = \int_0^1 |\psi(\{u + \{k_n \gamma\}\}) - \psi(u)|^2 du \rightarrow 0.$$

Nun ist andererseits

$$\psi(\{u + k\gamma\}) = e^{i\lambda(f(u) + f(\{u + \gamma\}) + \dots + f(\{u + (k-1)\gamma\}))} \cdot \psi(u),$$

also

$$\begin{aligned} & \int_0^1 |\psi(\{u + k\gamma\}) - \psi(u)|^2 du \\ &= \int_0^1 |(e^{i\lambda(f(u) + f(\{u + \gamma\}) + \dots + f(\{u + (k-1)\gamma\}))} - 1) \psi(u)|^2 du \\ &= c^2 \int_0^1 |e^{i\lambda(f(u) + f(\{u + \gamma\}) + \dots + f(\{u + (k-1)\gamma\}))} - 1|^2 du \\ &= 4c^2 \int_0^1 \sin^2 \left(\frac{\lambda}{2} (f(u) + f(\{u + \gamma\}) + \dots + f(\{u + (k-1)\gamma\})) \right) du. \end{aligned}$$

Wenn wir also

$$\liminf_{k \rightarrow \infty} \left[\int_0^1 \sin^2 \left(\frac{\lambda}{2} (f(u) + f(\{u + \gamma\}) + \dots + f(\{u + (k-1)\gamma\})) \right) du \right] > 0$$

⁶⁹ Es ist evident für alle $\chi(u) = e^{2\pi i k u}$, $k = 0, \pm 1, \pm 2, \dots$, also gilt es auch für deren Linearaggregate und für die Häufungspunkte der Linearaggregate — also für alle $\chi(u)$.

⁷⁰ Kronecker,

sowie Weyl, a. a. O. Anm. ⁵⁴.

beweisen können, so ist die Unerfüllbarkeit der $\psi(u)$ -Bedingung, d. h. die Nichtexistenz des Punktspektrums (abgesehen von $\lambda = 0$, $\varphi = \text{Konstanz}$) erwiesen.

4. Wir nehmen nun von $f(u)$, das wir periodisch mit der Periode 1 über alle u fortsetzen (so daß wir statt $f(\{u\})$ $f(u)$ schreiben können), folgendes an: Abgesehen von endlich vielen Punkten $x^{(1)}, \dots, x^{(\nu)}$ (sowie den ihnen mod der Periode 1 kongruenten Punkten $x^{(1)} + k, \dots, x^{(\nu)} + k$, $k = 0, \pm 1, \pm 2, \dots$; $x^{(1)}, \dots, x^{(\nu)}$ seien $\geq 0, < 1$) sei $f(u)$ überall stetig und stetig differentiierbar. In $x^{(1)}, \dots, x^{(\nu)}$ möge es die bzw. Sprünge $\alpha^{(1)}, \dots, \alpha^{(\nu)}$ machen. Die Summe dieser Sprünge, $\alpha = \alpha^{(1)} + \dots + \alpha^{(\nu)}$, sei $\neq 0$.

Wir teilen das Intervall $0 \leq u < 1$ in $(\nu+1)k$ Teile ein: I_l sei das Intervall $\frac{l-1}{(\nu+1)k} \leq u < \frac{l}{(\nu+1)k}$, $l = 1, 2, \dots, (\nu+1)k$. Die $\{\alpha^{(1)} - j\gamma\}, \dots, \{\alpha^{(\nu)} - j\gamma\}$, $j = 0, 1, \dots, k-1$ sind νk Zahlen, mindestens k unter den $(\nu+1)k$ Intervallen I_l sind also von ihnen frei, etwa: $I_{l^{(1)}}, \dots, I_{l^{(\nu)}}$. In diesen sind also $u, \{u+\gamma\}, \dots, \{u+(k-1)\gamma\}$ alle $\neq \alpha^{(1)}, \dots, \alpha^{(\nu)}$, d. h. $f(u) + f(u+\gamma) + \dots + f(u+(k-1)\gamma)$ stetig. Also ist in $I_h^{(u)}$, $h = 1, \dots, k$,

$$\begin{aligned} & f(u) + f(u+\gamma) + \dots + f(u+(k-1)\gamma) \\ &= C_h + \int_0^u (f'(x) + f'(x+\gamma) + \dots + f'(x+(k-1)\gamma)) du. \end{aligned}$$

Nun gilt für alle beschränkten, nach Riemann integrierbaren Funktionen $g(u)$ (mit der Periode 1)

$$\frac{g(u) + g(u+\gamma) + \dots + g(u+(k-1)\gamma)}{k} \rightarrow \int_0^1 g(x) dx,$$

und zwar gleichmäßig in u .⁷¹ Setzen wir für $g(x)$ $f'(x)$ ein und beachten

$$\int_0^1 f'(u) du + \alpha^{(1)} + \dots + \alpha^{(\nu)} = 0, \quad \int_0^1 f'(u) du = -\alpha^{(1)} - \dots - \alpha^{(\nu)} = -\alpha,$$

so sehen wir, daß

$$\frac{f'(u) + f'(u+\gamma) + \dots + f'(u+(k-1)\gamma)}{k} \rightarrow -\alpha \quad \text{für } k \rightarrow \infty$$

gleichmäßig in u gilt. Wegen $\alpha \neq 0$ hat also

$$\frac{f'(u) + f'(u+\gamma) + \dots + f'(u+(k-1)\gamma)}{k}$$

⁷¹ Weyl, a. a. O. Anm.⁵⁴. Der dort geführte Beweis der Konvergenz ergibt auch die gleichmäßige Konvergenz.

bei hinreichend großem k für alle u dasselbe Vorzeichen und einen Absolutwert $\geq \frac{1}{2}|\alpha|$, $\leq \frac{3}{2}|\alpha|$.

Also ist in $I_{l^{(n)}}$

$$f(u) + f(u + \gamma) + \dots + f(u + (k-1)\gamma) = C_h \pm \frac{|\alpha|k}{2} \int_0^u \sigma_h(x) dx,$$

wobei stets $1 \leq \sigma_h(x) \leq 3$ ist. Die Gesamtvariation dieses Ausdrucks in $I_{l^{(n)}}$, dessen Länge $\frac{1}{(\nu+1)k}$ ist, ist $\leq \frac{1}{(\nu+1)k} \cdot \frac{|\alpha|k}{2} \cdot 3 = \frac{3|\alpha|}{\nu+1}$, also $< \left| \frac{\pi}{\lambda} \right|$, wenn ν groß genug ist. Dies darf aber vorausgesetzt werden, denn wir können ν durch Hinzufügen fiktiver Unstetigkeitspunkte (mit Sprüngen = 0) beliebig erhöhen. Die Gesamtvariation ist also $<$ als die halbe Distanz von zwei benachbarten $\frac{2\pi n}{\lambda}$, $n = 0, \pm 1, \pm 2, \dots$. Ferner variiert der obige Ausdruck in ganz $I_{l^{(n)}}$ monoton. Mindestens in einer Hälfte von $I_{l^{(n)}}$ variiert er also (bei wachsendem u) vom nächsten $\frac{2\pi n}{\lambda}$ weg. Da er ferner im Viertel von $I_{l^{(n)}}$ um $\geq \frac{1}{4} \cdot \frac{1}{(\nu+1)k} \cdot \frac{|\alpha|k}{2} = \frac{|\alpha|}{8(\nu+1)}$ wächst, ist er mindestens in einem Viertel von $I_{l^{(n)}}$ um $\geq \frac{|\alpha|}{8(\nu+1)}$ vom nächsten $\frac{2\pi n}{\lambda}$ entfernt. Sein $\frac{\lambda}{2}$ -faches, das Argument im $\sin^2(\dots)$ des $\liminf_{k \rightarrow \infty} \left[\int_0^1 \sin^2(\dots) du \right]$ am Ende von 3., ist also dort um $\geq \frac{|\alpha| \cdot |\lambda|}{16(\nu+1)}$ vom nächsten πn entfernt. Da aber $\sin^2 y$ stets $\geq \left(\frac{2 \cdot \text{Distanz des nächsten } \pi n \text{ von } y}{\pi} \right)^2$ ist, ist der $\sin^2(\dots)$ selbst $\geq \frac{\alpha^2 \lambda^2}{64 \pi^2 (\nu+1)^2}$.

Im Integral $\int_0^1 \sin^2(\dots) du$ am Ende von 3. sind also, falls k groß genug ist, k Intervalle (die genannten Viertel der $I_{l^{(k)}}$, $h = 1, \dots, k$) vorhanden, jedes von der Länge $\frac{1}{4(\nu+1)k}$, in deren jedem der Integrand $\geq \frac{\alpha^2 \lambda^2}{64 \pi^2 (\nu+1)^2}$ ist. Da der Integrand stets ≥ 0 ist, ist das ganze Integral jedenfalls

$$\geq k \cdot \frac{1}{4(\nu+1)k} \cdot \frac{\alpha^2 \lambda^2}{64 \pi^2 (\nu+1)^2} = \frac{\alpha^2 \lambda^2}{256 \pi^2 (\nu+1)^2} = \delta > 0,$$

also

$$\liminf_{k \rightarrow \infty} \left[\int_0^1 \sin^2(\dots) du \right] \geq \delta > 0.$$

Damit ist das am Ende von 3. Gesagte bewiesen.

5. Wir haben also gezeigt:

SATZ 6. Sei S_t die in 1. mit Hilfe der Funktion $f(u)$ (die periodisch mit der Periode 1 fortgesetzt zu denken ist) und der Zahl γ definierte Strömung. $f(u)$ sei überall stetig und stetig differentierbar, abgesehen von endlich vielen Stellen in $0 \leq u < 1$ (sowie den diesen mod der Periode 1 kongruenten Stellen), wo Unstetigkeiten in Form von Sprüngen (d. i. $f(u+0)$, $f(u-0)$ existieren, aber $f(u+0) - f(u-0) \neq 0$) zulässig sind. Die Derivierte sei beschränkt und nach Riemann integrierbar. Ferner sei $0 \leq \gamma < 1$.

Dann ist für den (eigentlich-) ergodischen Charakter die Irrationalität von γ notwendig und hinreichend; für die Nichtexistenz eines Punktspektrums (außer $\lambda = 0$, $\varphi = \text{Konstanz}$) das Nichtverschwinden der Summe der Sprünge von $f(u)$ (in $0 \leq u < 1$) hinreichend.

Das reine Streckenspektrum ist also hier der allgemeine Fall.

Wie weit unsere Bedingung notwendig ist, ist schwerer zu entscheiden. Wenn $f(u)$ überall stetig ist, oder die Summe seiner Sprünge Null ist, kann ein Punktspektrum auftreten. Z. B.:

Sei $f(u)$ ein trigonometrisches Polynom:

$$f(u) = a_0 + \sum_1^n (a_m \cos 2\pi m u + b_m \sin 2\pi m u).$$

(Wegen $A \leq f(u) \leq B$ ist $A \leq a_0 \leq B$). Damit $\psi(\{u+\gamma\}) = e^{i\lambda f(u)} \psi(u)$ lösbar sei, genügt die Lösbarkeit von $\chi(\{u+\gamma\}) - \chi(u) = f(u) \cdots \bmod \frac{2\pi}{\lambda}$, da wir $\psi(u) = e^{i\lambda \chi(u)}$ setzen können. Da aber

$$\chi_1(\{u+\gamma\}) - \chi_1(u) = \sum_1^n (a_m \cos 2\pi m u + b_m \sin 2\pi m u)$$

lösbar ist⁷², genügt es $\chi(\{u+\gamma\}) - \chi(u) = a_0 \cdots \bmod \frac{2\pi}{\lambda}$ zu lösen. Setzen wir $\chi(u) = \alpha u$ an, so ist

$$\chi(\{u+\gamma\}) - \chi(u) = \alpha \gamma \text{ bzw. } = \alpha(\gamma - 1),$$

also muß dann

$$\alpha \gamma - a_0 = \frac{2\pi k}{\lambda}, \quad \alpha \gamma - \alpha - a_0 = \frac{2\pi l}{\lambda} \quad (k, l = 0, \pm 1, \pm 2, \dots)$$

sein, also

⁷² Es genügt $\chi_1(\{u+\gamma\}) - \chi_1(u) = \cos 2\pi m u$ und $= \sin 2\pi m u$ zu lösen, d. h. $= \Re e^{2\pi i m u}$ und $= \Im e^{2\pi i m u}$. Dies leistet

$$\chi_1(u) = \Re \text{ bzw. } \Im \frac{e^{2\pi i m u}}{e^{2\pi i m \gamma} - 1}.$$

$$\alpha = \frac{2\pi(k-l)}{\lambda}, \text{ und } a_0 = \alpha\gamma - \frac{2\pi k}{\lambda} = \frac{2\pi\gamma(k-l) - 2\pi k}{\lambda},$$

$$\lambda = \frac{2\pi\gamma}{a_0}(k-l) - \frac{2\pi}{a_0}k.$$

Wenn wir also statt $k-l$, $-k$ k , l schreiben, sehen wir: die

$$\lambda = \frac{2\pi\gamma}{a_0}k + \frac{2\pi}{a_0}l \quad (k, l = 0, \pm 1, \pm 2, \dots)$$

gehören alle zum Punktspektrum.

Oder es sei $f(u) \begin{cases} = 1 & \text{für } 0 \leq u < \frac{1}{2} \\ = 2 & \text{für } \frac{1}{2} \leq u < 1 \end{cases}$. (Die Sprünge sind: -1 bei $u = 0$, 1 bei $u = \frac{1}{2}$, ihre Summe ist 0.) Denn wenn $\lambda = 2\pi k$ ($k = 0, \pm 1, \pm 2, \dots$) ist, ist $e^{i\lambda f(u)} = 1$, also $\psi(u) = \text{Konstanz}$ Lösung.

Andererseits gibt es stetige $f(u)$, für die ein reines Streckenspektrum vorliegt, z. B.

$$f(u) = a_0 + a_1 \sum_1^\infty \frac{\cos(2^{(2^n)} \pi u)}{2^{(2^n)}} \quad \left(a_0 > |a_1| \sum_1^\infty \frac{1}{2^{(2^n)}} \right),$$

falls in der dyadischen Entwicklung von γ für unendlich viele n die $2^{(2^n)}$ te Ziffer eine 1 ist, und alle weiteren Ziffern, bis zur $2^{(2^{n+1})}$ ten (ausschl.) 0ten. Man beweist dies durch explizite Ausrechnung von $f(u) + f(u+\gamma) + \dots + f(u+(k-1)\gamma)$ für $k = 2^{(2^m)}$, und Abschätzung des Einflusses der einzelnen Glieder von $\sum_1^\infty$, wobei dasjenige mit $n = m$ dominiert — dann

zeigt sich, daß der $\liminf_{m \rightarrow \infty} \left[\int_0^1 \sin^2(\dots) du \right]$ vom Ende von 3. > 0 ist, obwohl $\{2^{(2^m)}\gamma\}$ nach Annahme beliebig kleine Werte annimmt. Wir übergehen die Einzelheiten dieser Rechnung. Vermutlich liegt bei den meisten stetigen $f(u)$ ein reines Streckenspektrum vor.

Unsere Strömung S_t läßt sich noch auf anschaulichere Formen bringen (durch isomorphe Abbildungen), doch soll dies bei einer anderen Gelegenheit erörtert werden.

ZUSÄTZE ZUR ARBEIT „ZUR OPERATORENMETHODE ...“¹².

1. Als erstes soll ein Fehler berichtigt werden, der dem Verf. beim Abfassen der zitierten Arbeit unterlaufen ist. Im Hauptsatz von Teil II, § 3, d. i. Satz 2 (der die Zerlegbarkeit aller Strömungen in ergodische Strömungen aussagt), wird u. a. behauptet, daß die dort konstruierten Mengen $X^{(0)}$ und $M(x)$ m -Räume sind³, also u. a. vollständig³. Dabei wurde übersehen, daß die zu Satz 2 führende Konstruktion die Vollständigkeit nicht sicherstellt. Die Vollständigkeit ist aber insofern wichtig, als z. B. der zweite Teil von Satz 5⁴ auf Satz 1 der vorhergehenden Abhandlung des Verf. beruht⁵, der seinerseits die Vollständigkeit als Prämisse hat.

Da jeder absolute G_δ -Raum bekanntlich so ummetrisiert werden kann, daß seine Topologie ungeändert bleibt und er vollständig wird⁶, würde es genügen, $X^{(0)}$ und die $M(x)$ als G_δ -Mengen zu wählen, um diese Lücke auszufüllen — was in der Tat durchführbar ist. Wir ziehen aber vor, zu zeigen, daß die Vollständigkeit für unsere Zwecke, d. h. für die Gültigkeit von Satz 1, 2 und 3, a. a. O. Anm.⁵, entbehrlich ist — d. h. daß der Begriff des „ m -Raumes“ ohne sie⁷ gefaßt werden kann. Dies ist erwiesen, wenn es uns gelingt, einen metrisch separablen Raum Ω , in dem ein l -Maß μ^* definiert ist, derart zu einem metrischen, separablen und vollständigen Raume $\bar{\Omega}$ mit einem l -Maß $\bar{\mu}^*$ zu erweitern, daß $\bar{\mu}^*$ für Mengen $\subset \Omega$ mit μ^* übereinstimmt und $\bar{\mu}^*(\bar{\Omega} - \Omega) = 0$ ist. (Denn dann überträgt sich die Prämisse von Satz 1 ohne weiteres von Ω , Ω' auf $\bar{\Omega}$, $\bar{\Omega}'$ und seine Aussage von $\bar{\Omega}$, $\bar{\Omega}'$ auf Ω , Ω' . Dasselbe gilt für Satz 2 ebendort, allerdings ist dort die Existenz einer eindeutigen Abbildung von $\bar{\Omega} - \Omega$ auf $\bar{\Omega}' - \Omega'$ notwendig, was z. B. sicher geht, wenn beide die Mächtigkeit des Kontinuums haben⁸. Satz 3 gilt auch, da er direkt aus Satz 2 folgt.)

¹ Annals of Math., Bd. 33, 3. Heft (1932), S. 587–642.

² Received September 6, 1932.

³ Vgl. die Definition D_1 auf S. 588 a. a. O.

⁴ S. 629 a. a. O.

⁵ Annals of Math., Bd. 33, 3. Heft (1932), S. 574–586; vgl. insbesondere Def. 1–5 und Satz 5.

⁶ Vgl. z. B. Hausdorff, „Mengenlehre“ (Berlin und Leipzig, 1927), S. 214, Satz III.

⁷ D. h. als metrisch und separabel sowie, wie sich w. u. zeigen wird, absolut Borelsch.

⁸ Sei C die nirgendsdichte Cantorsche perfekte Menge in ihrer ursprünglichen Topologie, $\bar{\mu}^*$ werde für alle Teilmengen von C gleich 0 definiert. Wir fügen $\bar{\Omega}$, C zu $\bar{\Omega}$ zusammen, u. zw. unter Beibehaltung ihrer Topologien und $\bar{\mu}^*$, aber voneinander isoliert. Dann spielt $\bar{\Omega}$ die Rolle von $\bar{\Omega}$, und $\bar{\Omega} - \Omega \supset C$ hat die Mächtigkeit des Kontinuums.

Wählen wir $\bar{\Omega}$ als vollständige Hülle von Ω^9 und definieren für die $M \subset \bar{\Omega}$ $\bar{\mu}^*(M) = \mu^*(M \cdot \Omega)$ ($M \cdot \Omega \subset \Omega$), so ist alles Gewünschte erfüllt, bloß bezüglich des l -Maß-Charakters des neuen $\bar{\mu}^*$ sind I., V.¹⁰ zu verifizieren. Dies gelingt in der Tat:

Ad I.: Ist K eine Kugel in $\bar{\Omega}$ mit einem Mittelpunkt aus Ω , so ist $K \cdot \Omega$ eine Kugel in Ω , also $\bar{\mu}^*(K) = \mu^*(K \cdot \Omega)$ endlich; gehört der Mittelpunkt von K nicht zu Ω , so ist doch K Teilmenge einer solchen Kugel, also $\bar{\mu}^*(K)$ wieder endlich.

Ad V.: Dies braucht nach Anm.⁹ a. a. O. nur mit Borelschen O bewiesen zu werden. Wenn Ω absolut Borelsch ist, d. h. Borelsch in seiner vollständigen Hülle $\bar{\Omega}^{11}$, so schließen wir so: Zu $M \subset \bar{\Omega}$, also $M \cdot \Omega \subset \Omega$, gibt es ein in Ω , also auch in $\bar{\Omega}$ Borelsches $N \supset M \cdot \Omega$ mit $\mu^*(N) = \mu^*(M \cdot \Omega)$. Dann ist $N + (\bar{\Omega} - \Omega)$ auch Borelsch, $\supset M$, und sein $\bar{\mu}^*$ gleich $\mu^*(N) = \mu^*(M \cdot \Omega) = \bar{\mu}^*(M)$, was zu beweisen war.

Wir haben also gezeigt:

Die wesentlichen Sätze 1, 2 und 3 aus der in Anm.⁵ zitierten Arbeit gelten auch, wenn der Begriff des m -Raumes weiter gefaßt wird: indem neben der Metrizität und Separabilität keine Vollständigkeit, sondern bloß absolut Borelscher Charakter¹¹ gefordert wird.

Zu Satz 2 der in Anm.¹ zitierten Arbeit ist dann zu sagen: da die $M(x)$ (in Ω) Borelsch sind, sind sie in diesem Sinne gleichzeitig mit Ω m -Räume ($X^{(0)}$ kann durch seinen maßgleichen F_σ -Kern ersetzt werden, es ist also auch Borelsch, d. h. m -Raum). Im neuen Sinne gilt also Satz 2 ohne weiteres.

2. Im Zusammenhange mit dem Problemkreise des Ergodensatzes sind noch zwei Abhandlungen von Herrn T. Carleman zu nennen, die nach Abfassung der w. o. zitierten Arbeiten des Verf. erschienen sind:

- I. Application de la théorie des équations intégrales linéaires aux équations différentielles de la dynamique, Arkiv för Mat., Astr. och Fys., 22 B (1932), No. 7, S. 1–5, eingegangen am 27. Mai 1931, erschienen am 2. März 1932;
- II. Application ... aux systèmes d'équations différentielles non linéaires, Acta Math., 59 (1932), S. 63–87, eingegangen am 10. Februar 1932, erscheint demnächst.

Carleman gelangt in diesen Abhandlungen unabhängig zu Koopmans Operatorenansatz und zu einem, mit demjenigen des Verf. verwandten Beweise des Ergodensatzes¹² unter Beschränkung auf differentiiierbare Strömungen.

⁹ Vgl. Hausdorff, S. 106–107.

¹⁰ Aus Def. 2, a. a. O., Anm. ⁵.

¹¹ Vgl. Hausdorff, S. 121 und 208.

¹² Vgl. die Arbeit Koopmans, Proc. Nat. Ac., Bd. 17, Mai 1931, S. 315; diejenige des Verf., Proc. Nat. Ac., Bd. 18, Jan. 1932, S. 70; sowie Teil II, § 2 der in Anm.¹ zitierten.

Da für diese der zur Schar U_t gehörige Operator R (vgl. zur Bezeichnung die Arbeit des Verf. a. a. O. Anm. ¹, S. 592, und die zweitvorhergehende, S. 571, sowie a. a. O. Anm. ¹², S. 71) direkt angebbar und seine Hypermaximalität leicht beweisbar ist, kann der Stonesche Satz vermieden und unmittelbar die Theorie des Spektrums von R benutzt werden. Diesen Weg schlägt Carleman ein.

Indessen ist sein Beweisverfahren in I. wesentlich lückenhaft, denn seine Betrachtungen beruhen auf der dortigen Formel (4), welche bei den meisten Strömungen unrichtig ist. In der Tat zeigt der Vergleich von (4) mit der vom Verf. benutzten Form der Spektraltheorie, daß das dort auftretende $\theta(p_0, q_0 | \lambda)$ der Integralkern des Operators $1 - E\left(\frac{1}{\lambda}\right)$ ($\lambda > 0$) bzw. $-E\left(\frac{1}{\lambda}\right)$ ($\lambda < 0$) sein muß. Bereits auf der Energiefläche des harmonischen linearen Oscillators besitzen aber die genannten Operatoren überhaupt keinen Integralkern. (In Carlemans Ausdrucksweise ist dann $n = 1$, $0 \leq x_1 \leq 1$, $x_1 = 0$ mit $x_1 = 1$ zu identifizieren, $A_1 = 1$ — d. h. $\Omega(U) = i \frac{dU}{dx_1}$, falls die Einheiten geeignet gewählt werden. Es liegt also ein reines Punktspektrum bei $\frac{1}{\lambda} = 0, \pm 2\pi, \pm 4\pi, \dots$ vor, zu 0 gehört die Eigenfunktion 1. Somit wäre $\theta\left(p_0, q_0 \left| \frac{1}{\pi} \right.\right) - \theta\left(p_0, q_0 \left| -\frac{1}{\pi} \right.\right) + 1$ der Integralkern von $(1 - E(\pi)) - (-E(-\pi)) + P_{[1]} = 1 + P_{[1]} - (E(\pi) - E(-\pi)) = 1$, aber der Einheitsoperator 1 hat keinen Integralkern. Hierbei haben wir die Normierung $\theta(p_0, q_0 | 0) = 0$ wesentlich benutzt, aber andere Beispiele setzen auch diese nicht voraus.)

In II. wird dieser Fehler durch einen besonderen Kunstgriff (Einführung der Hilfsvariablen ξ, η an Stelle von p_0, q_0) vermieden. Die Rolle der Formel (4) von I. übernimmt die richtige Formel (19) von II., und der Beweis wird lückenlos.

DIE EINFÜHRUNG ANALYTISCHER PARAMETER IN TOPOLOGISCHEN GRUPPEN

Einleitung.

1. Hilbert formulierte 1900 die folgende Frage¹: Sei G eine durch n Parameter stetig beschriebene Gruppe von stetigen Transformationen einer m -parametrischen Mannigfaltigkeit M , ist es dann möglich, in G und M solche Parameter einzuführen, daß alle die die Transformationen beschreibenden Funktionen stetig differentiierbar, oder sogar analytisch², ausfallen? Anders ausgedrückt: ist auf solche G , bei geeigneter Parameterwahl, die Liesche Theorie anwendbar? (Vgl. a. a. O. Anm. ¹, ².)

Brouwer gab 1909/10 durch schwierige und scharfsinnige topologische Untersuchungen³ alle G , M -Typen für $m = 1, 2$ an, wodurch die Hilbertsche Frage für $m = 1$ bejahend beantwortet und für $m = 2$ weitgehend gefördert wurde.

Ein weiter unten anzugebendes Gegenbeispiel lehrt, daß die Hilbertsche Frage bei (in M) intransitivem G im allgemeinen zu verneinen ist. (Für $m = 1$ zeigt Brouwers Liste, daß nur transitive Gruppen existieren; unser Gegenbeispiel entspricht dem einfachsten der übrigen Fälle: $n = 1, m = 2$.) Wir werden daher im folgenden G stets als (in M) transitiv voraussetzen.

Unter dieser Voraussetzung werden wir für alle geschlossenen (in der topologischen Terminologie: kompakten) G zeigen, daß die Hilbertsche Frage zu bejahen ist. Für nichtkompakte G liegen vorläufig keine ähnlicherweise abschließenden Resultate vor.

Die vorliegenden Untersuchungen wurden durch die in Anm. ¹³ angeführte Arbeit von A. Haar angeregt, und sie sind auf die Resultate derselben aufgebaut. Der Verfasser möchte Herrn Haar an dieser Stelle seinen Dank aussprechen, daß er ihm dieselben noch vor der Publikation zur Verfügung stellte.

2. Ein wichtiger Spezialfall des Hilbertschen Problems ist $G = M$. Denn wenn wir die Elemente von G mit a, b, x, y bezeichnen und das

¹ Kongreßvortrag, gehalten am internationalen Mathematikerkongreß in Paris, 1900. Abgedruckt in den Gött. Nachr., 1900. Es handelt sich hier um das „Problem 5“.

² F. Schur bewies, Math. Ann. 41 (1893), S. 509—538, daß für alle jenen G , die in M transitiv sind, und für die eine Parameterwahl möglich ist, bei der alle genannten Funktionen zweimal stetig differentiierbar sind, auch eine solche möglich ist, bei der sie alle analytisch sind.

³ L. E. J. Brouwer, Math. Ann. 67 (1909), S. 246—267, und 69 (1910), S. 181—203.

Kompositionsgesetz in G mit $a \cdot b$, so bestimmt jedes feste Element a von G eine Transformation von G : $x \rightarrow x \cdot a^4$. Die Frage lautet dann: Ist es möglich, in G n derartige Parameter einzuführen, daß die das Kompositionsgesetz beschreibenden Funktionen alle stetig differentierbar bzw. analytisch ausfallen? (Natürlich sind stets Parameter „im kleiner“, d. h. in einer geeignet gewählten Umgebung der Einheit von G , bzw. eines beliebig gegebenen Punktes von M , gemeint.) Wir werden als erstes dieses Problem erledigen und dann das allgemeinere (beliebiges M) darauf zurückführen.

Es genügt natürlich zu zeigen, daß G auf eine Gruppe G' mit der genannten Eigenschaft homöomorph⁵ und isomorph⁶ abgebildet werden kann.

Eine Menge G' von k -dimensionalen komplexen Matrizen von nichtverschwindender Determinante ($k = 1, 2, \dots$), die die Einheitsmatrix E_k enthält, und mit den Matrizen A, B auch die Inverse A^{-1} und ihr Matrizenprodukt $A \cdot B$, ist eine Gruppe mit der Kompositionsregel $A \cdot B$. Sie kann topologisiert werden, indem wir die Matrizen A als Punkte des k^2 -dimensionalen komplexen, d. h. des $2k^2$ -dimensionalen reellen, euklidischen Raumes ansehen (vgl. a. a. O. Anm. ⁷), S. 5—6). Wir nehmen ferner an, daß jeder Häufungspunkt (d. i. Häufungsmatrix) von G' mit nichtverschwindender Determinante auch zu G' gehört. Dann kann G' in einer geeigneten Umgebung der Einheit homöomorph auf $h (= 0, 1, \dots, 2k^2)$ reelle Parameter $\alpha_1, \dots, \alpha_h$ (die in einer gewissen Umgebung von $0, \dots, 0$ variieren) bezogen werden⁷, ja die Matricelemente der Matrizen von G' sind sogar analytische Funktionen von $\alpha_1, \dots, \alpha_h$. Infolgedessen ist das Kompositionsgesetz von G' (d. i. das Multiplikationsgesetz der Matrizen, das in den Matricelementen algebraisch ist) in den Parametern $\alpha_1, \dots, \alpha_h$ analytisch. D. h. ein solches G' hat die von uns gewünschten Eigenschaften.

Ist nun G geschlossen (d. h. kompakt), und besitzt es eine stetige Darstellung $a \rightarrow D(a)$ durch k -dimensionale Matrizen, die auch treu (d. h. ein-eindeutig) ist, so können wir so schließen: Die Menge G' aller $D(a)$ (a durchläuft G) ist wegen der Kompaktheit von G abgeschlossen, und auch die inverse Abbildung ist stetig, d. h. die Darstellung ist eine homöomorphe Abbildung. Daß sie isomorph ist, ist klar. Da unsere Behauptung nach obigem für G' gilt⁸, gilt sie somit auch für G .

⁴ Übt man zuerst $x \rightarrow x \cdot a$ aus und dann $x \rightarrow x \cdot b$, so wird $x \rightarrow x \cdot (a \cdot b)$. G ist offenbar transitiv in sich selbst.

⁵ D. h. ein-eindeutig und gleichzeitig mit der Umkehrung stetig.

⁶ D. h. das Kompositionsgesetz von G in dasjenige von G' überführend.

⁷ Dies bewies der Verf., Math. Zeitschr. 30 (1929), S. 3—42. Vgl. insbesondere Satz I auf S. 27.

⁸ Eigentlich fehlt noch $\det(D(a)) \neq 0$. Wegen $D(a) \cdot D(a^{-1}) = D(1)$ steht dies fest, wenn $D(1)$ die Einheitsmatrix ist. Dies ist aber stets erreichbar (vgl. z. B. a. a. O. Anm. ⁷, in Anm. ¹², S. 28).

Wir brauchen also nur die Existenz treuer stetiger Darstellungen für die geschlossenen G nachzuweisen. Dieser, auch an und für sich nicht uninteressante, Satz wird daher der Hauptgegenstand unserer Betrachtungen sein.

Die genannte Darstellung $D(a)$ wird übrigens aus lauter unitären Matrizen bestehen. Daher können wir den letztgenannten Satz, unter Berücksichtigung des in Anm. ⁷ zitierten Satzes, auch so formulieren: Die einzigen (im kleinen) n -parametrischen ($n = 0, 1, 2, \dots$) geschlossenen (d. h. kompakten) Gruppen sind die abgeschlossenen Gruppen unitärer Matrizen des k -dimensionalen (komplexen) euklidischen Raumes ($k = 0, 1, 2, \dots$) sowie deren homöomorph-isomorphe Bilder.

3. Um den eigentümlichen, analytisch-gruppentheoretisch-topologischen Charakter der zu verwendenden Methode zu veranschaulichen, seien jene Resultate der neueren gruppentheoretischen und topologischen Literatur, auf der die zu führenden Beweise beruhen, hier zusammengestellt. Es sind die folgenden:

- a) A. Haar bewies vor kurzem, daß in jeder topologischen Gruppe⁹, welche separabel¹⁰ und im kleinen kompakt¹¹ ist, ein rechtsinvariantes Lebesguesches äußeres Maß¹² definiert werden kann¹³. Infolgedessen kann nach bekanntem Vorgang auch eine rechtsinvariante Lebesguesche Integration $\int_G f(x) dv_x$ (dv_x symbolisiert das „Volumelement“ für die Variable x) definiert werden¹⁴. Hierbei ist zu erwähnen, daß bereits F. Peter und H. Weyl, unter Vervollkommnung von Ansätzen von Hurwitz, eine solche Integration für solche n -parametrische Gruppen konstruierten, deren Kompositionsgesetz auf stetig differenzierbare Funktionen der Parameter führt¹⁵; das Wichtige an Haars Konstruktion ist aber gerade, daß sie ganz allgemein und rein mengentheoretisch, ohne jede

⁹ D. h. eine Gruppe, in der ein Umgebungsbegriff definiert ist, derart, daß die Komposition $a \cdot b$ und die Inverse a^{-1} stetige Funktionen von a, b bzw. von a sind. Vgl. für die Grundbegriffe der Topologie etwa Hausdorff, „Mengenlehre“, Berlin und Leipzig (1927), § 40, S. 226—230.

¹⁰ Vgl. a. a. O. Anm. ⁹, S. 125, und (10) auf S. 229.

¹¹ D. h. jeder Punkt hat eine Umgebung, deren abgeschlossene Hülle kompakt ist. Da die Abbildung $x \rightarrow x \cdot a$ homöomorph ist, genügt es (bei Gruppen), dies für die Einheit zu verlangen.

¹² Vgl. Carathéodory, „Reelle Funktionen“, Berlin und Leipzig (1918), S. 237—274; oder (für beliebige topologische Räume) in der Arbeit des Verf., Ann. of Math., 33 (1932), S. 574—586, Def. 2, 4 auf S. 574, 576. „Rechtsinvarianz“ besagt, daß die Abbildung $x \rightarrow x \cdot a$ maßtreu ist.

¹³ A. Haar, Ann. of Math., 34 (1933), S. 147—169, insbesondere § 3.

¹⁴ Z. B. als Stieltjessches Integral $\int_{-\infty}^{+\infty} \lambda d(\text{Maß}_x[f(x) \leq \lambda])$.

¹⁵ Math. Ann. 97 (1928), S. 737—755, insbesondere S. 737.

- Beziehung zu einer Parameterdarstellung, direkt zu einem Maßbegriff führt. Unsere Überlegungen, die vom Haarschen Maß zu den analytischen Parametern führen, können also als eine Art Umkehrung der Peter-Weylschen angesehen werden, bei denen das Umgekehrte geleistet wird — wenn auch die Schwierigkeiten von anderer Art sind.
- b) Bei geschlossenem (kompaktem) G ist das Haarsche Maß von ganz G , wie man leicht einsieht, endlich, kann also $= 1$ normiert werden. Als dann sind die Operatoren- und Eigenwertmethoden und -sätze anwendbar, die Peter und Weyl a. a. O. Anm. ¹⁵ verwenden, wodurch eine vollständige Theorie der (stetigen) Darstellungen von G durch unitäre endlich vieldimensionale Matrizen entsteht.
- c) Bis hierher wurde bloß benutzt, daß G eine topologische Gruppe ist. Unter Benutzung der Tatsache, daß es (im kleinen) homöomorph auf n Parameter bezogen werden kann, wird durch topologische Betrachtungen gezeigt, daß unter den in b) genannten Darstellungen auch treue vorkommen müssen. Hierbei werden in erster Linie zwei „dimensionstheoretische“ Sätze benutzt: Brouwers Satz von der Gebietsinvarianz und Lebesgues „Pflastersteinsatz“ bzw. der damit zusammenhängende „Allgemeine Zerlegungssatz“ der Dimensionstheorie von K. Menger und Urysohn ¹⁶.
- d) Bei den topologischen Betrachtungen werden auch die in Anm. ⁷ zitierten Resultate des Verf. über die Parameterdarstellung von Gruppen unitärer Matrizen mehrfach wesentlich herangezogen.

Die Einteilung der Arbeit ist diese: In Teil I wird die Existenz der treuen stetigen Darstellung bewiesen, in Teil II dies auf das Hilbertsche Problem angewendet und das Gegenbeispiel für (in M) intransitive G angegeben.

I. Einführung analytischer Parameter in n -parametrischen geschlossenen Gruppen G .

1. G sei eine topologische Gruppe, d. h. eine, in der ein Umgebungsbegriff definiert ist, derart, daß $a \cdot b$ und a^{-1} in a, b bzw. a stetig sind (vgl. Anm. ⁹). Wir nennen G n -parametrisch ($n = 1, 2, \dots$), wenn jedes a von G eine Umgebung $U(a)$ hat, die einer offenen Punktmenge $K(a)$ im n -dimensionalen (reellen) euklidischen Raume R_n (der Zahlen n -tupel $\alpha_1, \dots, \alpha_n$) homöomorph ist. Nach Anm. ⁹ genügt es, mit Hinblick auf die homöomorphe Abbildung $x \rightarrow x \cdot a$, diese Forderung nur für $a = 1$ (d. i. die Gruppeneinheit) zu stellen.

¹⁶ L. E. J. Brouwer, Math. Ann. **70** (1913), S. 161—165; H. Lebesgue, Fund. Math. **2** (1921), S. 256—285, sowie E. Sperner, Hamb. Abh. **6** (1928), S. 265—272; K. Menger, „Dimensionstheorie“, Berlin und Leipzig (1928), S. 251—266.

Aus demselben Grunde können wir alle $K(a)$ als einander gleich annehmen. Eine Translation von $K(a)$ im R_n erlaubt es, das Bild von a nach $0, \dots, 0$ zu verlegen; eine eventuelle Verkleinerung von $U(a)$ und $K(a)$, $K(a)$ als Kugel mit dem Mittelpunkt $0, \dots, 0$ anzunehmen; eine nachfolgende Ähnlichkeitstransformation in R_n , den Radius von $K(a)$ gleich 1 zu normieren. Also wird $K(a)$ gleich $K_1: \alpha_1^2 + \dots + \alpha_n^2 < 1$. Wir bilden noch die (abgeschlossene) Kugel $\bar{K}_1: \alpha_1^2 + \dots + \alpha_n^2 \leq \frac{1}{2}$. Bei der Abbildung $U(a) \rightarrow K(a) = K_1$ habe $\bar{K}_1$ das Urbild $\bar{V}(a)$.

$\bar{K}_1$ ist kompakt und separabel, also auch $\bar{V}(a)$; wenn also G durch endlich viele $\bar{V}(a)$ überdeckt werden kann, so ist es kompakt und separabel; wenn es durch abzählbar viele gelingt, ist es Vereinigung abzählbar vieler kompakter Mengen und separabel. Umgekehrt: da a innerer Punkt von $\bar{V}(a)$ ist, genügen zur Überdeckung von G , falls es kompakt ist, endlich viele $\bar{V}(a)$ und, falls es separabel ist, abzählbar viele. Wir haben also gezeigt:

G ist dann und nur dann kompakt bzw. separabel, wenn es durch endlich bzw. abzählbar viele $\bar{V}(a)$ überdeckt werden kann. Wenn es separabel ist, ist es allenfalls Vereinigung abzählbar vieler kompakter Mengen.

Wir werden im folgenden die Analytisierungsfrage für die „geschlossenen“, d. h. kompakten, G lösen, für die „offenen“, d. h. bloß separablen, G bleibt sie unerledigt. Im folgenden wird also G durchweg als n -parametrig und kompakt vorausgesetzt.

2. In G definierte stetige Funktionen, mit komplexen Zahlen als Werten, bezeichnen wir mit $f(x)$, $g(x)$, $k(x, y)$, $\varphi(x)$; die aufs Haarsche Maß gegründete Integration mit $\int_G f(x) dv_x$ (vgl. Einl. 3., insbesondere a) und Anm. ¹³, ¹⁴). Wie a. a. O. gesagt wurde, ist das Maß von G endlich, da dies für jede kompakte Menge der Fall ist; durch Multiplizieren aller Maße (und Integrale) mit dem konstanten Faktor $\frac{1}{\text{Maß } G}$ können wir daher erreichen, daß $\text{Maß } G = 1$ wird¹⁷. Nun können wir die in Einl. 3. b) zitierten Betrachtungen von Peter und Weyl auf unser G übertragen. (Vgl. auch Haar, a. a. O. Anm. ¹³, S. 167—169.)

Jedem stetigen Kern $k(x, y)$ ordnen wir den Integraloperator I_k zu, der durch

$$I_k f(x) = \int_G k(x, y) f(y) dv_y$$

¹⁷ Dies ist die einzige Stelle, wo die Kompaktheit von G wirklich wesentlich benötigt wird.

definiert ist. Die Methode von E. Schmidts Dissertation¹⁸ ist auf dieses I_k wörtlich anwendbar, wenn es die Hermitesche Symmetrie hat, d. h. wenn $k(y, x) = \overline{k(x, y)}$ gilt. Infolgedessen treffen insbesondere die folgenden Behauptungen zu:

Die von O verschiedenen Eigenwerte von I_k bilden eine Folge $\lambda_1, \lambda_2, \dots$ (die Anzahl ihrer Elemente ist $0, 1, 2, \dots$ oder abzählbar unendlich). Für jedes λ_ν gibt es unter den Eigenfunktionen, d. h. den Lösungen von $I_k f = \lambda_\nu f$, nur endlich viele linear unabhängige, ihre Maximalanzahl sei N_ν ($= 1, 2, \dots$). N_ν solche Eigenfunktionen können sogar normiert-orthogonal¹⁹ gewählt werden: $f_1^{(\nu)}, \dots, f_{N_\nu}^{(\nu)}$. Jede Funktion von der Form $I_k f$ (f beliebig) kann durch Linearaggregate endlich vieler $f_l^{(\nu)}$ ($\nu = 1, 2, \dots, l = 1, \dots, N_\nu$) in G gleichmäßig beliebig gut approximiert werden.

Sei nun $\varphi(x)$ eine stetige Funktion in G ; wir setzen $k(x, y) = \varphi(xy^{-1})$, $I_k = T_\varphi$, und wenden das soeben Gesagte auf T_φ an. Hierzu ist allerdings $\varphi(yx^{-1}) = \overline{\varphi(xy^{-1})}$ notwendig, d. h. $\varphi(x^{-1}) = \overline{\varphi(x)}$.

T_φ hat die folgende Eigentümlichkeit: wenn es $f(x)$ in $g(x)$ überführt, so führt es $f(x \cdot a)$ in $g(x \cdot a)$ über (man ersetze in der definierenden Formel für $I_k = T_\varphi$ die Integrationsvariable y durch $y \cdot a^{-1}$ und beachte die Rechtsinvarianz sowie $x \cdot (y \cdot a^{-1})^{-1} = (x \cdot a) \cdot y^{-1}$). Also ist mit $f(x)$ auch $f(x \cdot a)$ Eigenfunktion von T_φ , und zwar vom selben Eigenwert; also ist $f_l^{(\nu)}(x \cdot a)$ Linearaggregat der $f_1^{(\nu)}(x), \dots, f_{N_\nu}^{(\nu)}(x)$. Wir wollen jedes endliche normiert-orthogonale Funktionensystem $f_1(x), \dots, f_N(x)$, für welches jedes $f_l(x \cdot a)$ ($l = 1, \dots, N, a$ in G) Linearaggregat der $f_1(x), \dots, f_N(x)$ ist, ein Darstellungssystem nennen. Die $f_1^{(\nu)}(x), \dots, f_{N_\nu}^{(\nu)}(x)$ sind somit Darstellungssysteme, und wir erkennen: Es gibt eine Folge (von $0, 1, 2, \dots$ oder abzählbar unendlich vielen) Darstellungssystemen, so daß jede Funktion $T_\varphi f$ (φ fest, f beliebig) durch Linearaggregate endlich vieler ihrer Funktionen beliebig gut approximierbar ist.

Gäbe es eine Folge $\varphi_1, \varphi_2, \dots$ von φ 's, so daß jedes f durch $T_{\varphi_s} g$'s ($s = 1, 2, \dots, g$ beliebig) beliebig gut approximierbar ist, so erhielten wir durch Zusammenfassung der Darstellungssystemfolgen von $\varphi_1, \varphi_2, \dots$ eine neue Darstellungssystemfolge, derart, daß jedes f durch Linearaggregate endlich vieler ihrer Funktionen beliebig gut approximierbar ist. Hieraus folgt noch, daß diese Folge nicht leer und auch nicht endlich sein kann.

¹⁸ Math. Ann. 63 (1907), S. 433—476.

¹⁹ Ein Funktionensystem $f_1, f_2, \dots$ heißt normiert-orthogonal, wenn

$$\int_G f_j(x) \overline{f_l(x)} dv_x \begin{cases} = 1, & \text{für } j = l \\ = 0, & \text{für } j \neq l \end{cases}$$

Eine Folge $\varphi_1, \varphi_2, \dots$, für die $T_{\varphi_s} f$ für $s \rightarrow \infty$ gleichmäßig in G gegen f konvergiert (für jedes stetige f), hat die obige Eigenschaft. Wenn $U_1, U_2, \dots$ sich auf 1 (Gruppeneinheit von G) zusammenziehende Umgebungen von 1 sind, und $\varphi_s(x)$ stets reell und ≥ 0 , aber außerhalb von $U_s = 0$ ist und $\int_G \varphi_s(x) dv_x = 1$ gilt, sowie natürlich $\varphi(x^{-1}) = \varphi(x)$ ($= \overline{\varphi(x)}$), so leisten diese $\varphi_1, \varphi_2, \dots$ offenbar das Gewünschte. Solche φ_s sind aber zu den U_s leicht angebar: Zunächst genügt $\int_G \varphi_s(x) dv_x > 0$ (statt $= 1$), da wir sonst $\varphi_s(x)$ durch $\frac{\varphi_s(x)}{\int_G \varphi_s(x) dv_x}$ ersetzen. Ferner ist $\varphi_s(x^{-1}) = \varphi_s(x)$

unnötig, da wir sonst $\varphi_s(x)$ durch $\varphi_s(x) + \varphi_s(x^{-1})$ ersetzen, gleichzeitig U_s so weit verringernd, daß für alle x des neuen U_s x, x^{-1} beide zum alten gehören. Weiter können wir U_s in $\bar{V}(1)$ liegend annehmen, sein bei $\bar{V}(1) \rightarrow \bar{K}(\frac{1}{2})$ entstehendes Bild sei $\bar{U}$. $\bar{U}$ enthält gewiß eine Kugel mit dem Mittelpunkt $0, \dots, 0$, deren Radius sei $\epsilon_s > 0$. Liegt x in U_s , so heiße die Distanz seines Bildpunktes (in $\bar{K}(\frac{1}{2})$) von $0, \dots, 0$ $d_s(x)$. Dann erfüllt

$$\varphi_s(x) \left\{ \begin{array}{l} = \text{Max } (\epsilon_s - d_s(x), 0), \text{ für } x \text{ in } U_s \\ = 0, \text{ sonst} \end{array} \right|$$

die übriggebliebenen Bedingungen.

Somit existiert die oben beschriebene Folge von Darstellungssystemen, wir nennen sie $f_1^{(\mu)}(x), \dots, f_{M_\mu}^{(\mu)}(x)$ ($\mu = 1, 2, \dots$).

3. $f_l^{(\mu)}(x \cdot a)$ ist Linearaggregat der $f_1^{(\mu)}(x), \dots, f_{M_\mu}^{(\mu)}(x)$, also ist

$$f_l^{(\mu)}(x \cdot a) = \sum_{j=1}^{M_\mu} \alpha_{jl}^{(\mu)}(a) f_j^{(\mu)}(x) \quad (\mu = 1, 2, \dots, l = 1, \dots, M_\mu, a \text{ in } G).$$

Da die $f_j^{(\mu)}(x)$ normiert-orthogonal sind, ist

$$\alpha_{jl}^{(\mu)}(a) = \int_G f_l^{(\mu)}(x \cdot a) \overline{f_j^{(\mu)}(x)} dv_x,$$

so daß $\alpha_{jl}^{(\mu)}(a)$ stetig von a abhängt. Ferner ist die Matrix $A^{(\mu)}(a) = \{\alpha_{jl}^{(\mu)}(a)\}_{j,l=1,\dots,M_\mu}$ unitär, und es gilt offenbar $A^{(\mu)}(a \cdot b) = A^{(\mu)}(a) A^{(\mu)}(b)$. Die Matrizen $A^{(\mu)}(a)$ (μ fest) bilden also eine stetige Darstellung von G .

Wenn für a, b von G $A^{(\mu)}(a) = A^{(\mu)}(b)$ für jedes $\mu = 1, 2, \dots$ gilt, so ist nach Definition stets $f_l^{(\mu)}(x \cdot a) = f_l^{(\mu)}(x \cdot b)$. Daher muß $f(x \cdot a) = f(x \cdot b)$ für jedes f gelten, das Linearaggregat endlich vieler $f_l^{(\mu)}$ oder durch solche beliebig gut approximierbar ist — d. h. für jedes stetige f . $x = b^{-1}$ ergibt: für jedes stetige f ist $f(b^{-1}a) = f(1)$, also muß $b^{-1}a = 1$, $a = b$ sein²⁰.

²⁰ Ist $c \neq 1$, so sei U eine c nichtenthaltende Umgebung von 1. Konstruieren wir f zu U so, wie weiter oben φ_s zu U_s , so wird $f(1) > 0$, $f(c) = 0$.

D. h.: Die Gesamtheit der Darstellungen $A^{(1)}(a), A^{(2)}(a), \dots$ ist treu — aber wir brauchen eine einzelne (endlich vieldimensionale!) Darstellung, die es ist!

Sei $M_1 + \dots + M_\mu = L_\mu$, reihen wir $A^{(1)}(a), \dots, A^{(\mu)}(a)$ an der Diagonale einer L_μ -dimensionalen Matrix aneinander und besetzen die übrigen Felder mit Nullen, so entsteht eine L_μ -dimensionale Matrix $B^{(\mu)}(a)$ (vgl. Abb.); tun

Zeilennummern		$\leftarrow B^{(\infty)}(a) \rightarrow$			
		$\leftarrow B^{(1)}(a) \rightarrow$	$B^{(2)}(a) \rightarrow$	$\dots$	$B^{(\mu)}(a) \rightarrow \dots$
1 =	1				
2 =	2				
$\vdots$	$\vdots$				
$L_1 =$	M_1	$A^{(1)}(a)$	0	$\dots$	$\dots$
$L_1 + 1 =$	$M_1 + 1$				
$L_1 + 2 =$	$M_1 + 2$				
$\vdots$	$\vdots$				
$L_2 =$	$M_1 + M_2$	0	$A^{(2)}(a)$	$\ddots$	
$\vdots$	$\vdots$				
$L_{\mu-1} + 1 =$	$M_1 + \dots + M_{\mu-1} + 1$				
$L_{\mu-1} + 2 =$	$M_1 + \dots + M_{\mu-1} + 2$				
$\vdots$	$\vdots$				
$L_\mu =$	$M_1 + \dots + M_{\mu-1} + M_\mu$		0		$A^{(\mu)}(a)$
$\vdots$	$\vdots$				

ABILDUNG.

wir dasselbe mit allen $A^{(1)}(a), A^{(2)}(a), \dots$, so entsteht eine unendlich vieldimensionale, aber zeilen- und kolonnenunendliche Matrix $B^{(\infty)}(a)$ (vgl. Abb.). Die $B^{(\mu)}(a)$ sind offenbar unitäre Darstellungen von G , $B^{(\infty)}(a)$ ebenfalls. $B^{(\mu)}(a)$ hängt stetig von a ab, $B^{(\infty)}(a)$ ebenfalls, falls für unendliche Matrizen die Konvergenz durch die Konvergenz aller Matrizenelemente definiert wird²¹. Schließlich ist $B^{(\mu)}(a)$ der L_μ -te Abschnitt von $B^{(\infty)}(a)$. Die Menge aller $B^{(\mu)}(a)$ (a durchläuft ganz G) heiße $\mathfrak{G}^{(\mu)}$, diejenige aller $B^{(\infty)}(a)$ $\mathfrak{G}^{(\infty)}$; $\mathfrak{G}^{(\mu)}$ ist stetiges Bild von G , $\mathfrak{G}^{(\infty)}$ nach dem weiter oben Gesagten ein-eindeutiges und stetiges Bild — da G kompakt ist, muß daher auch die Inverse der letzteren Abbildung stetig sein.

Im folgenden werden wir zeigen, daß schon ein $\mathfrak{G}^{(\mu)}$ ein-eindeutiges Bild von G , d. h. schon ein $B^{(\mu)}(a)$ treu ist — womit das Programm von Einl. 2 erfüllt sein wird. Dazu werden wir aber die bisher kaum benutzte Beziehung von G zum n -dimensionalen (reellen) euklidischen Raume, d. i. die

²¹ Diese Topologie der unendlichen Matrizen $\{\alpha_{ji}\}_{j,i=1,2,\dots}$ kann aus verschiedenen Entfernungsbegriffen (vgl. a. a. O. Anm. ⁹, S. 94) hergeleitet werden, z. B. aus

$$\text{Entf.}(\{\alpha_{ji}\}, \{\beta_{ji}\}) = \min_{L=1,2,\dots} \left\{ \max_{j,i=1,\dots,L} \left[\max_{j,i=1,\dots,L} (|\alpha_{ji} - \beta_{ji}|), \frac{1}{L} \right] \right\}.$$

Abbildung von $U(1)$ auf $K(1)$, wirklich heranziehen müssen. (Bisher benutzten wir eigentlich nur, daß G topologisch, separabel und kompakt ist.)

4. Da $B^{(\mu)}(a)$ eine Darstellung ist, bildet es die Gruppe G auf eine Gruppe ab, d. h. $\mathfrak{G}^{(\mu)}$ ist eine Matrizengruppe. Wegen $1^2 = 1$ ist das Bild von 1 das eigene Quadrat, und da es unitär ist, muß es die L_μ -dimensionale Einheitsmatrix sein: E_{L_μ} . Da $B^{(\mu)}(a)$ stetig und G kompakt ist, ist auch das Bild kompakt, d. h. $\mathfrak{G}^{(\mu)}$ ist eine abgeschlossene Matrizenmenge. Somit sind die in Anm. ⁷ zitierten Resultate des Verf. auf $\mathfrak{G}^{(\mu)}$ anwendbar, insbesondere gilt folgendes (vgl. a. a. O., Satz I, auf S. 27):

Es gibt $P_\mu (= 0, 1, \dots, 2L_\mu^2)$ linear unabhängige, L_μ -dimensionale Matrizen $U_1^{(\mu)}, \dots, U_{P_\mu}^{(\mu)}$ (die Matrizenelemente sind komplex, die lineare Unabhängigkeit ist aber mit reellen Koeffizienten gemeint), so daß alle Matrizen $\exp(\alpha_1 U_1^{(\mu)} + \dots + \alpha_{P_\mu} U_{P_\mu}^{(\mu)})^{22}$ zu $\mathfrak{G}^{(\mu)}$ gehören. Werden die $\alpha_1, \dots, \alpha_{P_\mu}$ auf eine geeignete Umgebung $K^{(\mu)}$ des Punktes $0, \dots, 0$ beschränkt, so erschöpfen die

$$\exp(\alpha_1 U_1^{(\mu)} + \dots + \alpha_{P_\mu} U_{P_\mu}^{(\mu)})$$

sogar eine Umgebung $\mathfrak{R}^{(\mu)}$ der Einheitsmatrix E_{L_μ} in $\mathfrak{G}^{(\mu)}$ ein-eindeutig.

Fassen wir ein $U_p^{(\mu)}$ ins Auge. Sei $\varepsilon_1, \varepsilon_2, \dots$ eine gegen 0 konvergierende Folge von Zahlen > 0 . $\exp(\varepsilon_s U_p^{(\mu)})$ gehört zu $\mathfrak{G}^{(\mu)}$, es ist also wenigstens einem $B^{(\mu)}(a)$ gleich, etwa für $a = a_p$. Die $B^{(\mu+1)}(a_p)$ liegen in $\mathfrak{G}^{(\mu+1)}$, sie haben daher einen Häufungspunkt. Wenn wir $\varepsilon_1, \varepsilon_2, \dots$ durch eine Teilfolge ersetzen, können wir also erreichen, daß die $B^{(\mu+1)}(a_p)$ konvergieren. Auch das monotone Abnehmen der $\varepsilon_1, \varepsilon_2, \dots$ ist durch Aussondern einer Teilfolge erreichbar. Wir dürfen daher annehmen, daß beides für die ursprünglichen $\varepsilon_1, \varepsilon_2, \dots$ der Fall war. Es ist

$$\varepsilon_s - \varepsilon_{s+1} \rightarrow 0, \quad \varepsilon_s - \varepsilon_{s+1} > 0,$$

und

$$\begin{aligned} B^{(\mu)}(a_s a_{s+1}^{-1}) &= B^{(\mu)}(a_s) B^{(\mu)}(a_{s+1})^{-1} \\ &= \exp(\varepsilon_s U_p^{(\mu)}) \exp(\varepsilon_{s+1} U_p^{(\mu)}) = \exp((\varepsilon_s - \varepsilon_{s+1}) U_p^{(\mu)}), \\ B^{(\mu+1)}(a_s a_{s+1}^{-1}) &= B^{(\mu+1)}(a_s) B^{(\mu+1)}(a_s)^{-1} \rightarrow E_{L_{\mu+1}}, \end{aligned}$$

wenn wir also $\varepsilon_1, \varepsilon_2, \dots$ durch $\varepsilon_1 - \varepsilon_2, \varepsilon_2 - \varepsilon_3, \dots$ ersetzen, so tritt $a'_s = a_s a_{s+1}^{-1}$ mit $B^{(\mu+1)}(a'_s) \rightarrow E_{L_{\mu+1}}$ an die Stelle von a_s . Wir dürfen daher annehmen, daß $B^{(\mu+1)}(a_s) \rightarrow E_{L_{\mu+1}}$ schon für die ursprünglichen $\varepsilon_1, \varepsilon_2, \dots$ der Fall war.

²² Ist A eine endlich vieldimensionale Matrix, so ist $\exp(A)$ als $\lim_{h \rightarrow \infty} \sum_{s=0}^h \frac{A^s}{s!}$ definiert. Vgl. a. a. O. Anm. ⁷, S. 7–15.

Fast alle $B^{(\mu+1)}(a_s)$, $B^{(\mu)}(a_s)$ gehören also zu $\mathfrak{K}^{(\mu+1)}$ bzw. $\mathfrak{K}^{(\mu)}$, für diese ist

$$B^{(\mu+1)}(a_s) = \exp(\alpha_{s,1} U_1^{(\mu+1)} + \dots + \alpha_{s,P^{(\mu+1)}} U_{P^{(\mu+1)}}^{(\mu+1)}), \\ \alpha_{s,1}, \dots, \alpha_{s,P^{(\mu+1)}} \text{ aus } K^{(\mu+1)}.$$

Für $s \rightarrow \infty$ strebt dieser Ausdruck gegen $E_{L_{\mu+1}}$, also $\alpha_{s,1}, \dots, \alpha_{s,P_{\mu+1}}$ gegen $0, \dots, 0$, also strebt

$$\alpha_s = |\alpha_{s,1}| + \dots + |\alpha_{s,P_{\mu+1}}| > 0$$

(wäre es $= 0$, so wäre $\alpha_{s,1} = \dots = \alpha_{s,P_{\mu+1}} = 0$, $B^{(\mu+1)}(a_s) = E_{L_{\mu+1}}$, also erst recht $\exp(\epsilon_s U_p^{(\mu)}) = B^{(\mu)}(a_s) = E_{L_\mu}$, also $\epsilon_s = 0$, entgegen der Annahme) auch gegen 0 . Sei

$$\beta_{s,1} = \frac{\alpha_{s,1}}{\alpha_s}, \dots, \beta_{s,P_{\mu+1}} = \frac{\alpha_{s,P_{\mu+1}}}{\alpha_s},$$

dann ist

$$|\beta_{s,1}| + \dots + |\beta_{s,P_{\mu+1}}| = 1.$$

Für eine Teilfolge konvergieren die $\beta_{s,1}, \dots, \beta_{s,P_{\mu+1}}$; wir dürfen daher annehmen, daß dies schon für die ursprünglichen $\epsilon_1, \epsilon_2, \dots$ der Fall war. Die Limites seien bzw. $\bar{\beta}_1, \dots, \bar{\beta}_{P_{\mu+1}}$, es ist

$$|\bar{\beta}_1| + \dots + |\bar{\beta}_{P_{\mu+1}}| = 1.$$

Nennen wir den L_μ -Abschnitt von $U_q^{(\mu+1)}$ $\tilde{U}_q^{(\mu+1)}$. Da $B^{(\mu)}(a_s)$ der L_μ -Abschnitt von $B^{(\mu+1)}(a_s)$ ist, d. h. $\exp(\epsilon_s U_p^{(\mu)})$ der von

$$\exp(\alpha_s(\beta_{s,1} U_1^{(\mu+1)} + \dots + \beta_{s,P_{\mu+1}} U_{P_{\mu+1}}^{(\mu+1)})),$$

ist

$$\exp(\epsilon_s U_p^{(\mu)}) = \exp(\alpha_s(\beta_{s,1} \tilde{U}_1^{(\mu+1)} + \dots + \beta_{s,P_{\mu+1}} \tilde{U}_{P_{\mu+1}}^{(\mu+1)})).$$

Sobald ϵ_s, α_s klein genug sind, d. h. für fast alle s , folgt hieraus

$$\epsilon_s U_p^{(\mu)} = \alpha_s(\beta_{s,1} \tilde{U}_1^{(\mu+1)} + \dots + \beta_{s,P_{\mu+1}} \tilde{U}_{P_{\mu+1}}^{(\mu+1)}).$$

Somit konvergiert für $s \rightarrow \infty$

$$\frac{\epsilon_s}{\alpha_s} U_p^{(\mu)} \text{ gegen } \bar{\beta}_1 \tilde{U}_1^{(\mu+1)} + \dots + \bar{\beta}_{P_{\mu+1}} \tilde{U}_{P_{\mu+1}}^{(\mu+1)}.$$

Da sowohl $U_p^{(\mu)}$ als auch der Limes $\neq 0$ ist, folgt hieraus, daß $\frac{\epsilon_s}{\alpha_s}$ gegen einen Limes $\gamma \neq 0$ strebt; dann muß aber

$$U_p^{(\mu)} = \frac{\bar{\beta}_1}{\gamma} \tilde{U}_1^{(\mu+1)} + \dots + \frac{\bar{\beta}_{P_{\mu+1}}}{\gamma} \tilde{U}_{P_{\mu+1}}^{(\mu+1)}$$

sein, d. h. $U_p^{(\mu)}$ ist der L_μ -Abschnitt von $\frac{\bar{\beta}_1}{\gamma} U_1^{(\mu+1)} + \dots + \frac{\bar{\beta}_{P_{\mu+1}}}{\gamma} U_{P_{\mu+1}}^{(\mu+1)}$.

Somit ist jedes $U_p^{(\mu)}$, $p = 1, \dots, P_\mu$, L_μ -Abschnitt eines Linearaggregats der $U_q^{(\mu+1)}$, $q = 1, \dots, P_{\mu+1}$. Diese Linearaggregate müssen linear unabhängig sein, weil es schon ihre L_μ -Abschnitte $U_p^{(\mu)}$ sind, also ist $P_\mu \leq P_{\mu+1}$. Ferner können wir die $U_q^{(\mu+1)}$, $q = 1, \dots, P_{\mu+1}$, so umbezeichnen (nämlich linear transformieren), daß diese P_μ Linearaggregate gerade bzw. den $U_1^{(\mu+1)}, \dots, U_{P_\mu}^{(\mu+1)}$ gleich werden. Nehmen wir diese Transformation nacheinander für $\mu = 1, 2, \dots$ vor, so erreichen wir allgemein: $U_p^{(\mu)}$ ist der L_μ -Abschnitt von $U_p^{(\mu+1)}$. Dies werde im folgenden vorausgesetzt.

5. Wir betrachten für jedes $\nu \geq \mu \exp(x_1 U_1^{(\nu)} + \dots + x_{P_\mu} U_{P_\mu}^{(\nu)})$ und nennen eines derjenigen a , deren $B^{(\nu)}(a)$ diesem Ausdruck gleich ist, $a_\nu(x_1 \dots x_{P_\mu})$. Die $a_\nu(x_1 \dots x_{P_\mu})$ haben für $\nu \rightarrow \infty$ einen Häufungspunkt $a(x_1 \dots x_{P_\mu})$, weil G kompakt ist. Ferner ist für $\nu \geq \mu_1 (\geq \mu)$ der L_{μ_1} -Abschnitt von $B^{(\infty)}(a_\nu(x_1 \dots x_{P_\mu}))$ derselbe wie derjenige von

$$B^{(\nu)}(a_\nu(x_1 \dots x_{P_\mu})) = \exp(x_1 U_1^{(\nu)} + \dots + x_{P_\mu} U_{P_\mu}^{(\nu)}),$$

und zwar gleich

$$B^{(\mu_1)}(a_{\mu_1}(x_1 \dots x_{P_\mu})) = \exp(x_1 U_1^{(\mu)} + \dots + x_{P_\mu} U_{P_\mu}^{(\mu)}),$$

also von ν unabhängig. Daher ist dies auch der L_{μ_1} -Abschnitt von $B^{(\infty)}(a(x_1 \dots x_{P_\mu}))$, d. h. fast alle $B^{(\infty)}(a_\nu(x_1 \dots x_{P_\mu}))$ haben denselben L_{μ_1} -Abschnitt wie $B^{(\infty)}(a(x_1 \dots x_{P_\mu}))$. Wegen $L_{\mu_1} \rightarrow \infty$ bedeutet das, daß $B^{(\infty)}(a_\nu(x_1 \dots x_{P_\mu}))$ für $\nu \rightarrow \infty$ gegen $B^{(\infty)}(a(x_1 \dots x_{P_\mu}))$ konvergiert (vgl. das Ende von 3. sowie Anm. ²¹), also auch, daß $a_\nu(x_1 \dots x_{P_\mu})$ gegen $a(x_1 \dots x_P)$ konvergiert (vgl. ebendort). Da der L_{μ_1} -Abschnitt von $B^{(\infty)}(a(x_1 \dots x_{P_\mu}))$

$$\exp(x_1 U_1^{(\mu_1)} + \dots + x_{P_\mu} U_{P_\mu}^{(\mu_1)})$$

ist, also stetig in $x_1, \dots, x_{P_\mu}$, ist jedes Matricelement von $B^{(\infty)}(a(x_1 \dots x_{P_\mu}))$ stetig in $x_1, \dots, x_{P_\mu}$ also (vgl. das Ende von 3.) auch dieses selbst sowie $a(x_1 \dots x_{P_\mu})$. Liegt $x_1, \dots, x_{P_\mu}$ in $K^{(\mu)}$, so hängt schon der L_μ -Abschnitt von $B^{(\infty)}(a(x_1 \dots x_{P_\mu}))$ ein-eindeutig von $x_1, \dots, x_{P_\mu}$ ab, um so mehr dieses selbst sowie $a(x_1 \dots x_{P_\mu})$. Da jeder L_{μ_1} -Abschnitt von $B^{(\infty)}(a(0 \dots 0))$ gleich $E_{L_{\mu_1}}$ ist, ist dieses die Einheitsmatrix, also $a(0 \dots 0) = 1$. Wenn wir also $x_1, \dots, x_{P_\mu}$ auf eine hinreichend kleine Umgebung $\tilde{K}^{(\mu)}$ von $0, \dots, 0$ beschränken, so liegt $a(x_1 \dots x_{P_\mu})$ jedenfalls in $U(1)$ (in G).

Also bildet $a(x_1 \dots x_{P_\mu})$ ein Gebiet $\tilde{K}^{(\mu)}$ im P_μ -dimensionalen euklidischen Raume ein-eindeutig und stetig auf eine Teilmenge von $U(1)$ ab, das seiner-

seits ein-eindeutiges und stetiges Bild eines Gebietes $K(1)$ im n -dimensionalen euklidischen Raume ist. Aus Brouwers Satz von der Gebietsinvarianz (vgl. a. a. O. Anm. ¹⁶) folgt daher $P_\mu \leq n$. Also ist $P_1 \leq P_2 \leq \dots \leq n$, von einem gewissen $\bar{\mu}$ an sind also die P_μ konstant: $P_\mu = P_{\mu+1} = \dots = P \leq n$.

6. Sei $\mu \geq \bar{\mu}$. $\mathfrak{G}^{(\mu)}$ ist in einer Umgebung des Punktes E_{L_μ} mit dem Gebiet $K^{(\mu)}$ im P -dimensionalen Raume homöomorph, also P -dimensional; wegen der Gruppeneigenschaft gilt das für jeden Punkt, also ist es als topologischer Raum P -dimensional (vgl. Mengers in Anm. ¹⁶ zitiertes Werk, S. 79—80). Als kompakte P -dimensionale Menge im $2L_\mu^2$ -dimensionalen euklidischen Raume der L_μ -dimensionalen Matrizen, fällt es unter den „Allgemeinen Zerlegungssatz“ der Dimensionstheorie (a. a. O. S. 156): Zu jedem $\varepsilon > 0$ können endlich viele Mengen $\mathfrak{S}_1, \dots, \mathfrak{S}_t$ gefunden werden (natürlich von μ, ε abhängig), so daß 1. jedes $\mathfrak{S}_i$ einen Diameter $\leq \varepsilon$ hat²³, 2. je $P+2$ Mengen $\mathfrak{S}_i$ keinen gemeinsamen Punkt haben, 3. $\mathfrak{S}_1 + \dots + \mathfrak{S}_t = \mathfrak{G}^{(\mu)}$ ist. Nun sei $\mathfrak{S}'_i$ die Menge aller $B^{(\infty)}(a)$ mit $B^{(\mu)}(a)$ in $\mathfrak{S}_i$. Wieder gilt: 2'. je $P+2$ Mengen $\mathfrak{S}'_i$ haben keinen gemeinsamen Punkt, 3'. $\mathfrak{S}'_1 + \dots + \mathfrak{S}'_t = \mathfrak{G}^{(\infty)} = \text{Menge aller } B^{(\infty)}(a)$. Ferner ist die Entfernung von $B^{(\infty)}(a)$ und $B^{(\infty)}(b)$ im Sinne von Anm. ²¹ $\leq \text{Max} \left(\frac{1}{\mu}, d \right)$, wenn d die Entfernung ihrer L_μ -Abschnitte, d. h. von $B^{(\mu)}(a)$, $B^{(\mu)}(b)$ ist. Somit ist der Diameter von $\mathfrak{S}'_i \leq \text{Max} \left(\frac{1}{\mu}, \bar{d}_i \right)$, wenn $\bar{d}_i$ der Diameter von $\mathfrak{S}_i$ ist, d. h. $\leq \text{Max} \left(\frac{1}{\mu}, \varepsilon \right)$. Für hinreichend großes μ ($\mu \geq \frac{1}{\varepsilon}$) gilt also auch: 1'. jedes $\mathfrak{S}'_i$ hat einen Diameter $\leq \varepsilon$. Nach der „Umkehrung des allgemeinen Zerlegungssatzes“ (a. a. O. S. 157) ist daher $\mathfrak{G}^{(\infty)} \leq P$ -dimensional. Da es aber G homöomorph ist, ist es, wie dieses, n -dimensional. Also ist $P \geq n$.

Somit ist $P = n$, d. h. $P_\mu = P_{\mu+1} = \dots = n$.

7. Die in 5. aufgestellte homöomorphe Abbildung von $x_1, \dots, x_{P_\mu}$ auf $a(x_1, \dots, x_{P_\mu})$ bildet nach dem soeben Bewiesenen für $\mu \geq \bar{\mu}$ (wir halten μ fest) die $x_1, \dots, x_n$ auf die $a(x_1 \dots x_n)$ ab, u. zw. ein Gebiet $\tilde{K}^{(\mu)}$ des n -dimensionalen euklidischen Raumes auf eine Teilmenge U von $U(1)$, das einem Gebiet des n -dimensionalen euklidischen Raumes homöomorph ist. Aus Brouwers Satz von der Gebietsinvarianz (vgl. a. a. O. Anm. ¹⁶) folgt daher, daß dieses Bild U innere Punkte besitzt.

²³ Um besseren Anschluß an die in Anm. ²¹ definierte Entfernung bei unendlich viel-dimensionalen Matrizen zu wahren, definieren wir bei P -dimensionalen Matrizen die Entfernung so (das ist der in a. a. O. Anm. ⁷ verwendeten, und zur gewöhnlichen Topologie führenden, topologisch gleichwertig):

$$\text{Entf.}(\{\alpha_{ji}\}, \{\beta_{ji}\}) = \text{Max}_{j, l=1, \dots, P} (|\alpha_{jl} - \beta_{jl}|).$$

Seien a, b Punkte von U , $B^{(\mu)}(a) = B^{(\mu)}(b)$ ($\mu \geq \bar{\mu}$). Es ist

$$a = a(x_1 \cdots x_n), \quad B^{(\mu)}(a) = \exp(x_1 U_1^{(\mu)} + \cdots + x_n U_n^{(\mu)}),$$

ebenso

$$b = b(y_1 \cdots y_n), \quad B^{(\mu)}(b) = \exp(y_1 U_1^{(\mu)} + \cdots + y_n U_n^{(\mu)}),$$

also

$$\exp(x_1 U_1^{(\mu)} + \cdots + x_n U_n^{(\mu)}) = \exp(y_1 U_1^{(\mu)} + \cdots + y_n U_n^{(\mu)}),$$

und da $x_1, \dots, x_n$ und $y_1, \dots, y_n$ zu $\tilde{K}^{(\mu)}$ gehören,

$$x_1 = y_1, \dots, x_n = y_n.$$

Also ist $a = b$. Für $a \neq b$ ist somit

$$B^{(\mu)}(a) \neq B^{(\mu)}(b).$$

Die Menge der $x \cdot c$, x in U , heie $U \cdot c$. Ist $\bar{a}$ innerer Punkt von U , so ist stets a innerer Punkt von $U \cdot (\bar{a}^{-1} \cdot a)$, d. h. jeder Punkt von G ist innerer Punkt eines $U \cdot c$. Da G kompakt ist, gengen endlich viele $U \cdot c$, etwa $U \cdot c_1, \dots, U \cdot c_{\bar{r}}$, um ganz G zu berdecken. Da auch in $U \cdot c$ aus $a \neq b$ $B^{(\mu)}(a) \neq B^{(\mu)}(b)$ folgt (denn $a \cdot c^{-1}, b \cdot c^{-1}$ gehren zu U , also ist wegen $a \cdot c^{-1} \neq b \cdot c^{-1}$ $B^{(\mu)}(a \cdot c^{-1}) \neq B^{(\mu)}(b \cdot c^{-1})$, $B^{(\mu)}(a) \cdot B^{(\mu)}(c)^{-1} \neq B^{(\mu)}(b) \cdot B^{(\mu)}(c)^{-1}$, $B^{(\mu)}(a) \neq B^{(\mu)}(b)$), erkennen wir: in jedem $U \cdot c_r$ ($r = 1, \dots, \bar{r}$) kommt hchstens ein a mit $B^{(\mu)}(a) = E_{L_\mu}$ vor, insgesamt gibt es also $\leq \bar{r}$ solche a .

Nennen wir die Menge der a mit $B^{(\mu)}(a) = E_{L_\mu}$ $\mathfrak{E}_\mu$, sie ist offenbar eine Untergruppe von G , $\mathfrak{E}_1 \supseteq \mathfrak{E}_2 \supseteq \dots$. Nach obigem ist sie fr $\mu \geq \bar{\mu}$ endlich. Die Elementzahlen der endlichen $\mathfrak{E}_\mu \supseteq \mathfrak{E}_{\mu+1} \supseteq \dots$ sind monoton-nichtwachsend, also von einer gewissen Stelle $\mu = \bar{\mu}$ an konstant. Die $\mathfrak{E}_\mu \supseteq \mathfrak{E}_{\mu+1} \supseteq \dots$ haben alle gleich viele Elemente, also ist $\mathfrak{E}_\mu = \mathfrak{E}_{\mu+1} = \dots$. Wenn a zu $\mathfrak{E}_\mu$ gehrt, so gehrt es demnach zu allen $\mathfrak{E}_\mu$, $\mu \geq \bar{\mu}$, mithin sind alle $B^{(\mu)}(a) = E_{L_\mu}$ (fr $\mu \geq \bar{\mu}$, also fr alle μ). Daher ist auch $B^{(\infty)}(a)$ die Einheitsmatrix, also $= B^{(\infty)}(1)$, aber $B^{(\infty)}(x)$ ist ein-eindeutig (vgl. 3.), also $a = 1$. Da somit $\mathfrak{E}_{\bar{\mu}}$ nur die 1 enthlt, ist fr $a \neq 1$ $B^{(\bar{\mu})}(a) \neq E_{L_{\bar{\mu}}}$, also fr $a \neq b$ $B^{(\bar{\mu})}(a) \neq B^{(\bar{\mu})}(b)$. Also ist $B^{(\bar{\mu})}(a)$ ein-eindeutig, d. h. treu.

Hieraus folgt, da $B^{(\bar{\mu})}(a)$ eine homomorphe Abbildung ist, G ist $\mathfrak{G}^{(\bar{\mu})}$ homomorph und isomorph, so da $B^{(\bar{\mu})}(a)$ und $\mathfrak{G}^{(\bar{\mu})}$ die Eigenschaften der in Einl. 2. errterten $D(a)$, G' besitzen. Wir formulieren das Resultat:

SATZ I. *Die einzigen topologischen Gruppen, die (im Sinne von I., 1.) n -parametrig und kompakt sind, sind die abgeschlossenen Gruppen aus unitren Matrizen endlich vieldimensionaler (komplexer²⁴) euklidischer Rume sowie die diesen homomorph-isomorphen Gruppen.*

²⁴Man kann sich auch aufs Reelle beschrnken: dazu mu die Dimensionszahl des betreffenden Raumes verdoppelt und die imaginre Einheit i durch die Matrix $\begin{Bmatrix} 0 & 1 \\ -1 & 0 \end{Bmatrix}$ ersetzt werden.

II. Transformationsgruppen G topologischer Räume M .

1. Eine topologische Gruppe G heie eine stetige Transformationsgruppe von M , wenn jedem a von G eine homomorphe (d. h. ein-eindeutige und stetige) Abbildung von M auf M zugeordnet ist: $F_a(\xi)$ (ξ durchluft M); dabei soll noch

1. die Funktionalgleichung

$$F_b(F_a(\xi)) = F_{a \cdot b}(\xi)$$

gelten,

2. $F_a(\xi)$ ist als Funktion von a (bei festgehaltenem ξ) stetig.

Wir setzen im folgenden G n -parametrig und kompakt sowie in M transitiv voraus, ber M machen wir dagegen keine weiteren Annahmen.

Somit ist G mit einer abgeschlossenen Gruppe unitrer Transformationen eines endlich viel- (etwa k -) dimensionalen Raumes homomorph und isomorph (Satz I), und wir knnen es vorlufig mit der betr. Gruppe identifizieren. Ferner zeichnen wir ein Element ξ_0 von M aus und bilden fr jedes ξ von M die Menge aller a von G mit $F_a(\xi_0) = \xi$: $H(\xi)$. $H(\xi)$ ist offenbar abgeschlossen und wegen der Transitivitt nicht leer; $H(\xi_0)$ hat die Gruppeneigenschaft, ist also eine abgeschlossene Untergruppe von G ; die $H(\xi)$ sind die Rechtsnebangruppen von $H(\xi_0)$. Jedes a von G gehrt genau einem $H(\xi)$ an: dem mit $\xi = F_a(\xi_0)$, und demnach hngt ξ stetig von a ab.

Wir wenden den in I., 4. fr die dortigen $\mathfrak{G}^{(\mu)}$ herangezogenen Satz auf G und $H(\xi_0)$ an, dadurch mgen die (k -dimensionalen) Matrizen $U'_1, \dots, U'_n$ und $U_1, \dots, U_l$ entstehen. Aus Brouwers Satz von der Invarianz der Dimensionszahl (a. a. O. Anm. ¹⁶) folgt, da G gleichzeitig n - und n' -parametrig ist, $n = n'$. Da $H(\xi_0)$ Untergruppe von G ist, sind die $U_1, \dots, U_l$ Linearaggregate der $U'_1, \dots, U'_n$. Somit ist $l \leq n$, und wir knnen die $U'_1, \dots, U'_n$ so umbezeichnen (nmlich linear transformieren), da

$$U'_1 = U_1, \dots, U'_l = U_l$$

wird. Wir lassen daher die Akzente berhaupt fort: zu G gehren $U_1, \dots, U_n$, zu $H(\xi_0)$ $U_1, \dots, U_l$ ($l \leq n$).

2. Wir betrachten die k -dimensionalen Matrizen $\exp(n_1 U_{l+1} + \dots + n_{n-l} U_n)$ ($n_1, \dots, n_{n-l}$ reell) und behaupten: es gibt eine Umgebung U^* von $0, \dots, 0$ im (reellen) $n - l$ -dimensionalen euklidischen Raume, so da fr $u_1, \dots, u_{n-l}$ und $v_1, \dots, v_{n-l}$ aus U^* und a, b aus $H(\xi_0)$

$\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n) \cdot a = \exp(v_1 U_{l+1} + \dots + v_{n-l} U_n) \cdot b$
die Folge

$$u_1 = v_1, \dots, u_{n-l} = v_{n-l}$$

(und damit auch $a = b$) hat. Letzteres kann auch

$$\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n)^{-1} \cdot \exp(v_1 U_{l+1} + \dots + v_{n-l} U_n) = a \cdot b^{-1}$$

geschrieben werden, d. h. es besagt, daß die linke Seite zu $H(\xi_0)$ gehört.

Wäre die obige Behauptung falsch, so gäbe es zwei Folgen $u_1^{(\nu)}, \dots, u_{n-l}^{(\nu)}$ und $v_1^{(\nu)}, \dots, v_{n-l}^{(\nu)}$ ($\nu = 1, 2, \dots$), beide gegen $0, \dots, 0$ konvergierend, so daß für kein ν

$$u_1^{(\nu)} = v_1^{(\nu)}, \dots, u_{n-l}^{(\nu)} = v_{n-l}^{(\nu)}$$

ist, und für jedes ν

$$\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n)^{-1} \exp(v_1 U_{l+1} + \dots + v_{n-l} U_n)$$

zu $H(\xi_0)$ gehört. Da dieser Ausdruck gegen E_k konvergiert, kann er nach dem in I., 4. verwendeten Satze für fast alle ν in der Form $\exp(w_1^{(\nu)} U_1 + \dots + w_l^{(\nu)} U_l)$ geschrieben werden, wobei auch $w_1^{(\nu)}, \dots, w_l^{(\nu)}$ gegen $0, \dots, 0$ konvergiert. Die vorliegende Gleichheit kann auch

$$\begin{aligned} \exp(w_1^{(\nu)} U_{l+1} + \dots + u_{n-l}^{(\nu)} U_n) \cdot \exp(w_1^{(\nu)} U_1 + \dots + w_l^{(\nu)} U_l) \\ = \exp(v_1^{(\nu)} U_{l+1} + \dots + v_{n-l}^{(\nu)} U_n) \end{aligned}$$

geschrieben werden, d. h. der Ausdruck

$$\exp(x_{l+1} U_{l+1} + \dots + x_n U_n) \cdot \exp(x_1 U_1 + \dots + x_l U_l)$$

hat bei Ersetzung der Variablenreihe $x_1, \dots, x_n$ durch $w_1^{(\nu)}, \dots, w_l^{(\nu)}, u_1^{(\nu)}, \dots, u_{n-l}^{(\nu)}$ und durch $0, \dots, 0, v_1^{(\nu)}, \dots, v_{n-l}^{(\nu)}$ denselben Wert. Wenn er, falls $x_1, \dots, x_n$ auf eine hinreichend kleine Umgebung von $0, \dots, 0$ beschränkt werden, $x_1, \dots, x_n$ eindeutig festlegt, so hat das Obige für fast alle ν $u_1^{(\nu)} = v_1^{(\nu)}, \dots, u_{n-l}^{(\nu)} = v_{n-l}^{(\nu)}$ zur Folge, d. h. einen Widerspruch mit der Annahme.

Wir entwickeln den obigen Ausdruck in eine (überall konvergente) Potenzreihe²⁵ und deuten alle Glieder von Graden ≥ 2 nur durch *** an. Er ist dann gleich:

$$1 + x_1 U_1 + \dots + x_n U_n + ***.$$

Die Matrizen $U_1, \dots, U_n$ sind (reell) linear unabhängig. Wir ergänzen sie zu einem vollständigen linear unabhängigen System, indem wir $U_{n+1}, \dots, U_{2k}$ irgendwie demgemäß hinzufügen. Den obigen Ausdruck minus 1 stellen wir dann als Linearaggregat von $U_1, \dots, U_{2k}$ dar, die Koeffizienten sind dann eindeutig bestimmt, u. zw. sind diejenigen von $U_1, \dots, U_n$, wenn wir alle Glieder von Graden ≥ 2 durch *** andeuten, gleich

$$x_1 + ***, \dots, x_n + ***.$$

²⁵ Die Potenzreihenbetrachtung ließe sich vermeiden, aber am Ende von 3. kommt ohnehin eine solche vor.

Wenn wir zeigen können, daß die Gleichungen

$$x_1 + *** = y_1, \dots, x_n + *** = y_n$$

nach $x_1, \dots, x_n$ (falls diese auf eine hinreichend kleine Umgebung von $0, \dots, 0$ beschränkt werden) aufgelöst werden können, so sind wir am Ziele. Sie können aber offenbar durch $y_1, \dots, y_n$ -Potenzreihen für $x_1, \dots, x_n$ aufgelöst werden. Damit ist die Existenz der Umgebung U^* von $0, \dots, 0$ mit der weiter oben behaupteten Eigenschaft bewiesen.

Diese Eigenschaft kann auch so formuliert werden: Wenn $u_1, \dots, u_{n-l}$ in U^* liegt, so liegt jedes $\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n)$ in einer anderen Menge $H(\xi)$. Das zugehörige $\xi = \xi(u_1, \dots, u_{n-l})$ in M ist somit eindeutige Funktion von $u_1, \dots, u_{n-l}$ und auch stetige Funktion derselben, da es stetig vom Ausdruck $\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n)$ abhängt.

3. In dem in I, 4. verwendeten Satze kam eine Umgebung $\mathfrak{R}^{(\omega)}$ von 1 in $\mathfrak{G}^{(\omega)}$ und eine Umgebung $K^{(\omega)}$ von $0, \dots, 0$ im P_μ -dimensionalen Raume vor; wir ersetzen jetzt $\mathfrak{G}^{(\omega)}$ durch $H(\xi_0)$, P_μ durch $n-l$, dann mögen $\mathfrak{R}^{(\omega)}$, $K^{(\omega)}$ $\mathfrak{S}$, H heißen. Wenn $u_1, \dots, u_{n-l}$ in U^* liegt und $w_1, \dots, w_l$ in H , so nimmt

$$a(u_1 \dots u_{n-l} w_1 \dots w_l) = \exp(u_1 U_{l+1} + \dots + u_{n-l} U_n) \cdot \exp(w_1 U_1 + \dots + w_l U_l)$$

denselben Wert nicht zweimal an: Denn dazu müßten die $\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n)$ im selben $H(\xi)$ liegen, also die $u_1, \dots, u_{n-l}$ dieselben sein; alsdann müßte $\exp(w_1 U_1 + \dots + w_l U_l)$ denselben Wert haben, also auch $w_1, \dots, w_l$ dieselben sein. Somit bildet es die beschriebene $u_1, \dots, u_{n-l} w_1, \dots, w_l$ -Menge homöomorph auf eine Teilmenge von G ab. Da aber die erstere Menge ein Gebiet des n -dimensionalen (reellen) euklidischen Raumes ist und da das Bild Teilmenge einer Umgebung der 1 in G ist, welche letztere einem ebensolchen Gebiet homöomorph angenommen werden kann (G ist n -parametrig!), muß nach dem Brouwerschen Satz von der Gebietsinvarianz (vgl. a. a. O. Anm.¹⁶, übrigens ist er hier, wo es sich durchweg um analytische Abbildungen handelt, selbstverständlich) das Bild ebenfalls offen sein (in G). Somit ist es eine Umgebung W der (dazugehörigen) 1.

Wenn ein ξ von M nahe genug bei ξ_0 liegt, so ist es ein $\xi(u_1 \dots u_{n-l})$, $u_1, \dots, u_{n-l}$ in U^* . Sonst gäbe es nämlich eine gegen ξ_0 konvergente Folge $\xi^{(\nu)}$ ($\nu = 1, 2, \dots$), in der das nie der Fall ist. Nun ist $\xi^{(\nu)}$ jedenfalls $= F_{a^{(\nu)}}(\xi_0)$, eine Teilfolge der $a^{(\nu)}$, konvergiert, da G kompakt ist; also können wir annehmen, daß es die $a^{(\nu)}$ -Folge selbst tut. Der Limes sei a_0 , dann ist $F_{a_0}(\xi_0)$ des Limes der $\xi^{(\nu)} = F_{a^{(\nu)}}(\xi_0)$, d. h. ξ_0 , so daß a_0 zu $H(\xi_0)$ gehört. Also ist auch $\xi^{(\nu)} = F_{a^{(\nu)} a_0^{-1}}(\xi_0)$, und $a^{(\nu)} a_0^{-1}$ konvergiert gegen 1; wir können daher annehmen, daß es schon die $a^{(\nu)}$ taten, d. h. $a_0 = 1$. Dann liegen fast alle $a^{(\nu)}$ in der Menge der

$$\exp(u_1 U_{l+1} + \dots + u_{n-l} U_n) \cdot \exp(w_1 U_1 + \dots + w_l U_l),$$

$u_1, \dots, u_{n-l}$ in U^* , $w_1, \dots, w_l$ in H . $\alpha^{(v)}$, also auch dieses

$$\exp(u_1 U_{l+1} + \dots + u_{n-l} U_{l+1}),$$

gehört zu $H(\xi^{(v)})$, daher ist

$$\xi^{(v)} = \xi(u_1, \dots, u_{n-l}),$$

entgegen der Voraussetzung.

$\xi(u_1, \dots, u_{n-l})$ bildet somit U^* homöomorph auf eine Teilmenge von M ab, von der ξ_0 innerer Punkt ist. Sei O eine in ihr enthaltene Umgebung von ξ_0 , wegen der Homöomorphie der Abbildung ist das Urbild U_0 von O eine (Relativ-) Umgebung von $0, \dots, 0$ in U^* — d. h. eine Umgebung von $0, \dots, 0$.

Damit ist die Umgebung O von ξ_0 in M homöomorph auf die $n-l$ Parameter $u_1, \dots, u_{n-l}$ in U_0 bezogen. In einer Umgebung $\mathfrak{R}$ von 1 in G ist, wie wir wissen, G durch $a = \exp(\alpha_1 U_1 + \dots + \alpha_n U_n)$ homöomorph auf die n Parameter $\alpha_1, \dots, \alpha_n$ in K bezogen. Das Kompositionsgesetz von G lautet:

$$\begin{aligned} \exp(\alpha_1 U_1 + \dots + \alpha_n U_n) \cdot \exp(\beta_1 U_1 + \dots + \beta_n U_n) \\ = \exp(\gamma_1 U_1 + \dots + \gamma_n U_n) \end{aligned}$$

(alles in K), u. zw. ist diese Gleichung in hinreichender Nähe von $0, \dots, 0$ eindeutig nach $\gamma_1, \dots, \gamma_n$ auflösbar. Das Transformationsgesetz von G in M lautet (soweit die linke Seite in W liegt):

$$\begin{aligned} \exp(u_1 U_{l+1} + \dots + u_{n-l} U_n) \cdot \exp(\alpha_1 U_1 + \dots + \alpha_n U_n) \\ = \exp(v_1 U_{l+1} + \dots + v_{n-l} U_n) \cdot \exp(w_1 U_1 + \dots + w_l U_l) \end{aligned}$$

($u_1, \dots, u_{n-l}$ und $v_1, \dots, v_{n-l}$ in U^* , $w_1, \dots, w_l$ in H), u. zw. ist auch diese Gleichung in hinreichender Nähe von $0, \dots, 0$ eindeutig nach $v_1, \dots, v_{n-l}$, $w_1, \dots, w_l$ auflösbar (es interessieren uns nur $v_1, \dots, v_{n-l}$). Wir haben noch zu zeigen, daß $\gamma_1, \dots, \gamma_n$ von $\alpha_1, \dots, \alpha_n$, $\beta_1, \dots, \beta_n$ und $v_1, \dots, v_{n-l}$, $w_1, \dots, w_l$ von $u_1, \dots, u_{n-l}$, $\alpha_1, \dots, \alpha_n$ analytische Funktionen sind. Die erste Aufgabe ist Spezialfall der zweiten: mit $l = 0$, $\alpha_1, \dots, \alpha_n$ für $u_1, \dots, u_n$, $\beta_1, \dots, \beta_n$ für $\alpha_1, \dots, \alpha_n$, $\gamma_1, \dots, \gamma_n$ für $v_1, \dots, v_n$; daher genügt es, diese zu betrachten. Wenn wir in der obigen Gleichung beide Seiten durch ihre (überall konvergenten) Potenzreihen ersetzen und alle Glieder von Graden ≥ 2 nur durch *** andeuten, so entsteht die Matrixgleichung:

$$\begin{aligned} 1 + \alpha_1 U_1 + \dots + \alpha_l U_l + (\alpha_{l+1} + u_1) U_{l+1} + \dots + (\alpha_n + u_{n-l}) U_n + \dots + *** \\ = 1 + w_1 U_1 + \dots + w_l U_l + v_1 U_{l+1} + \dots + v_{n-l} U_n + ***, \end{aligned}$$

woraus (da $U_1, \dots, U_n$ linear unabhängig sind) jedenfalls Zahlengleichungen

$$\begin{aligned}\alpha_1 + *** &= w_1 + ***, \dots, \alpha_l + \dots = w_l + ***, \\ \alpha_{l+1} + u_1 + *** &= v_1 + ***, \dots, \alpha_n + u_{n-l} + *** = v_{n-l} + ***\end{aligned}$$

folgen. Diese können offenbar durch $u_1, \dots, u_{n-l}, \alpha_1, \dots, \alpha_n$ -Potenzreihen für $v_1, \dots, v_{n-l}, w_1, \dots, w_l$ aufgelöst werden, womit die Analytizität dieser Größen bewiesen ist.

Unser Resultat lautet also:

SATZ II. Die topologische Gruppe G sei (im Sinne von I., 1.) n -parametrig und kompakt und (im Sinne von II., 1.) eine stetige Transformationsgruppe des topologischen Raumes M^{26} . Ferner sei G in M transitiv. Es ist möglich, eine Umgebung W von 1 in G und eine Umgebung U_0 eines beliebigen ξ_0 in M anzugeben, so daß W und U_0 homöomorph auf n bzw. m Parameter ($m = 0, 1, \dots, n^{27}$) bezogen werden können (die gewisse Umgebungen von $0, \dots, 0$ im n - bzw. m -dimensionalen [reellen] euklidischen Räume durchlaufen), so daß sowohl das Kompositionsgesetz von G , als auch das Transformationsgesetz von M durch G in diesen Parametern analytisch ausfällt.

4. Wir lassen jetzt die Annahme der Kompaktheit von G und seiner Transitivität in M fallen. Dagegen verschärfen wir die übrigen Annahmen: M sei in einer geeigneten Umgebung eines jeden seiner Punkte homöomorph auf m Parameter beziehbar, und $F_a(\xi)$ sei als Zweivariablenfunktion von a, ξ stetig²⁸. Es soll gezeigt werden, daß es u. U. unmöglich ist, die Parameter für a und ξ so zu wählen, daß $a \cdot b$ und $F_a(\xi)$ analytisch oder auch nur stetig differenzierbar ausfallen.

Es wird, wie in Einl. 1 erwähnt wurde, $n = 1, m = 2$ sein, u. zw. sei G die reelle Zahlengerade, d. h. a eine beliebige reelle Zahl, das Kompositionsgesetz die Addition $a + b$. M sei die reelle u, v -Ebene. $\varphi(r)$ sei eine für alle $r > 0$ definierte, stetige (reelle) Funktion, die zu a gehörige Transformation von M sei eine Drehung mit dem Mittelpunkt $0, 0$, aber um einen von r abhängigen Winkel auf jedem Kreise $u^2 + v^2 = r^2$: um $\varphi(r) a$. D. h.

$$\begin{aligned}F_a(u, v) &= u \cos \{ \varphi(\sqrt{u^2 + v^2}) a \} - v \sin \{ \varphi(\sqrt{u^2 + v^2}) a \}, \\ &u \sin \{ \varphi(\sqrt{u^2 + v^2}) a \} + v \cos \{ \varphi(\sqrt{u^2 + v^2}) a \}.\end{aligned}$$

Durch geeignete Wahl von $\varphi(r)$ werden wir die gewünschten Beispiele gewinnen.

²⁶ Man beachte, daß die Beziehung von M zu Parametern nicht Voraussetzung ist; sie wird vielmehr bewiesen!

²⁷ m tritt an Stelle des $n - l$ im Beweise.

²⁸ Im Falle der Kompaktheit und Transitivität wurde beides durch Satz I bewiesen, die a, ξ -Stetigkeit insbesondere dadurch, daß $F_a(\xi)$ sogar analytisch gewählt werden konnte.

Wenn $\varphi(1) = 0$ für alle $r > 0, \leq r_0$ ($r_0 > 0$, fest) gilt, aber für kein $r > r_0$, so liegt in jeder Umgebung des Punktes ξ_0 gleich r_0 , 0 ein Gebiet, in dem $F_a(\xi) = \xi$ ist, aber auch ein anderes Gebiet, wo dies nicht gilt. Da eine analytische Funktion ein solches Verhalten nicht zeigen kann, ist eine Parameterwahl, bei der $F_a(\xi)$ analytisch wird, unmöglich. Somit erledigt z. B. $\varphi(r) = \text{Max}(0, r - r_0)$ die Annahme der Analytizität.

Etwas schwieriger ist der Fall der stetigen Differenzierbarkeit. Hier schließen wir so, daß wir annehmen, daß in je einer Umgebung von 0 bzw. 0, 0 solche Parameter a' für a und u', v' für u, v eingeführt wurden, daß $a' \cdot b'$ und $F_{a'}(u', v')$ stetig differenzierbar sind, und daraus folgern, daß $\varphi(r)$ gewisse Eigenschaften besitzen muß. Alsdann wählen wir $\varphi(r)$ so, daß es dieselben nicht besitzt, womit auch die Annahme der stetigen Differenzierbarkeit entfällt.

5. Das Kompositionsgesetz für a' heiße $a' \circ b'$, die Abbildung von a auf a' werde durch die ein-eindeutige und stetige Funktion $a = f(a')$ vermittelt. $f(a')$ ist somit monoton (wachsend oder fallend), nach einem Satze von Lebesgue²⁹ gibt es also Stellen, wo es differenzierbar ist.

$a' \cdot b'$ ist als a' -Funktion differenzierbar, $a' \circ b'^{-1}$ auch; hätte also $a' \circ b'$ für ein $a' = a'_0$ die Ableitung 0, so hätte sie $(a' \circ b') \circ b'^{-1} = a'$ dort auch. Aber dies ist unmöglich, da a' stets die Ableitung 1 hat. Es ist

$$f(a' \circ b') = f(a') + f(b'),$$

sei $f(a')$ für $a' = a'_1$ differenzierbar; da $a' \circ b'$ die a' -Ableitung $\neq 0$ hat, ist auch $f(a')$ für $a' = a'_1 \cdot b'$ differenzierbar. D. h. $f(a')$ ist überall differenzierbar. Dabei ist

$$\frac{\partial}{\partial a'} f(a') \Big|_{a' = a'_1 \cdot b'} = \frac{\frac{\partial}{\partial a'} f(a') \Big|_{a' = a'_1}}{\frac{\partial}{\partial a'} (a' \circ b') \Big|_{a' = a'_1}},$$

also

$$\frac{\partial}{\partial a'} f(a') \Big|_{a' = \bar{a}'} = \frac{\frac{\partial}{\partial a'} f(a') \Big|_{a' = a'_1}}{\frac{\partial}{\partial a'} (a' \circ [\bar{a}' \circ a'^{-1}_1]) \Big|_{a' = a'_1}},$$

also $f(a')$ stetig differenzierbar. Da stets $\frac{\partial}{\partial a'} f(a') \neq 0$ ist (käme $\frac{\partial}{\partial a'} f(a') = 0$ überhaupt vor, so müßte es wegen der obigen Formel immer gelten, d. h. $f(a')$ wäre konstant), ist es auch die Inverse. Wir können daher statt a'

²⁹ Vgl. z. B. a. a. O. Anm. ¹², Satz 7 auf S. 573.

wieder den Parameter a einführen, $F_{a'}(u', v')$ bleibt stetig differenzierbar, d. h. wir können annehmen, daß von vornherein $a' = a$ war.

Wenn die zwei Komponenten von $\frac{\partial}{\partial a} F_a(u', v')$ ³⁰ für $a = \bar{a}$ und ein gewisses u', v' verschwinden, so müssen sie es wegen

$$F_{a+b}(u', v') = F_b(F_a(u', v'))$$

auch für $a = \bar{a} + b$ tun, d. h. für alle a . Dann ist $F_a(u', v')$ von a unabhängig, was $u, v = 0, 0$, also $u', v' = 0, 0$ nach sich zieht. Für $u', v' \neq 0, 0$ sind also die zwei Komponenten von $\frac{\partial}{\partial a} F_a(u', v')$ stets $\neq 0, 0$. $\varphi(\sqrt{u^2 + v^2})$ ist in u, v , also auch in u', v' stetig; wir nehmen an, daß es stets > 0 ist, dann ist auch $\frac{2\pi}{\varphi(\sqrt{u^2 + v^2})}$ in u', v' stetig. Da aber identisch

$$F \frac{2\pi}{\varphi(\sqrt{u^2 + v^2})}(u', v') = u', v'$$

gilt, so folgt aus dem soeben festgestellten Verhalten von $\frac{\partial}{\partial a} F_a(u, v)$, daß $\frac{2\pi}{\varphi(\sqrt{u^2 + v^2})}$ für $u', v' \neq 0, 0$ (nach u', v') stetig differenzierbar ist, also $\varphi(\sqrt{u^2 + v^2})$ auch.

Dem Kreis $u^2 + v^2 = r^2$ ($r > 0$, $< r_1$, wobei $u^2 + v^2 < r_1^2$ eine Umgebung von $0, 0$ ist, in der die Parameterdarstellung noch gilt) entspricht in der u', v' -Ebene eine Jordankurve J_r . Aus den Lagen der Kreise folgt, daß $0, 0$ und jedes J_s mit $s < r$ im Inneren von J_r liegt. $\varphi(\sqrt{u^2 + v^2})$ ist auf J_r konstant.

Sei C das Maximum von $\frac{\partial}{\partial u'} \varphi(\sqrt{u^2 + v^2})$ und $\frac{\partial}{\partial v'} \varphi(\sqrt{u^2 + v^2})$ im Inneren von J_{r_1} . Sei weiter $0 < r_2 < r_3 < \frac{r_1}{2}$, r_2, r_3 fest. Wir wählen ein u'_2, v'_2 auf J_{r_2} und ein u'_3, v'_3 auf J_{r_3} , ihre Verbindungsstrecke sei S . Seien nun $\nu - 1$ beliebige Zahlen $\varrho_1, \dots, \varrho_{\nu-1}$ mit $r_2 < \varrho_1 < \dots < \varrho_{\nu-1} < r_3$ gegeben, wir setzen noch $\varrho_0 = r_2, \varrho_\nu = r_3$. S schneidet jedes ϱ_μ ($\mu = 1, \dots, \nu - 1$), etwa in $u^{(\mu)}, v^{(\mu)}$, u. zw. können wir $u^{(\mu)}, v^{(\mu)}$ zwischen $u^{(\mu-1)}, v^{(\mu-1)}$ und u_3, v_3 wählen, was für alle $\mu = 2, 3, \dots, \nu$ geschehen soll. Naturgemäß

$$u^{(0)}, v^{(0)} = u'_2, v'_2, \quad u^{(\nu)}, v^{(\nu)} = u'_3, v'_3.$$

³⁰ $F_a(u', v')$ ist ein Punkt der u', v' -Ebene, besteht also aus zwei reellen Funktionen.

Dann liegen die $u^{(\mu)}, v^{(\mu)}$ in der Reihenfolge ihrer μ auf S , und es ist:

$$\begin{aligned} |\varphi(q_\mu) - \varphi(q_{\mu-1})| &\leq 2C \cdot \text{Entf.}(u'^{(\mu)}, v'^{(\mu)}; u'^{(\mu-1)}, v'^{(\mu-1)}), \\ \sum_{\mu=1}^{\nu} |\varphi(q_\mu) - \varphi(q_{\mu-1})| &\leq 2C \cdot \sum_{\mu=1}^{\nu} \text{Entf.}(u'^{(\mu)}, v'^{(\mu)}; u'^{(\mu-1)}, v'^{(\mu-1)}) \\ &= 2C \cdot \text{Entf.}(u'^{(\nu)}, v'^{(\nu)}; u'^{(0)}, v'^{(0)}) = C \cdot \text{Entf.}(u'_8, v'_8; u'_2, v'_2). \end{aligned}$$

Wir haben also, unter der Voraussetzung, daß für alle $r > 0, < r_1$ $\varphi(r) > 0$ ist, folgendes gezeigt: Es gibt zwei r_2, r_3 , $0 < r_2 < r_3 < \frac{r_1}{2}$, so daß für alle Systeme $q_0, q_1, \dots, q_\nu$ mit $r_2 = q_0 < q_1 < \dots < q_{\nu-1} < q_\nu = r_3$

$$\sum_{\mu=1}^{\nu} |\varphi(q_\mu) - \varphi(q_{\mu-1})|$$

unterhalb einer festen (d. h. nur von r_2, r_3 abhängigen) Schranke bleibt. Wenn wir also ein stetiges $\varphi(r)$ angeben können, derart, daß stets $\varphi(r) > 0$ ist und für jedes $r_2 < r_3$ die obere Grenze der

$$\sum_{\mu=1}^{\nu} |\varphi(q_\mu) - \varphi(q_{\mu-1})| \quad (r_2 = q_1 < q_2 < \dots < q_{\nu-1} < q_\nu = r_3)$$

gleich ∞ ist, so sind wir am Ziele. D. h. es wird ein stetiges positives $\varphi(r)$ gesucht, das in jedem Intervalle von unendlicher Variation ist, und Beispiele solcher $\varphi(r)$ sind bekannt³¹.

³¹ Vgl. z. B. a. a. O. Anm. ¹², S. 590–594; man bilde mit der dort konstruierten, „nirgends differenzierbaren Funktion von Weierstraß“ $w(x)$ $\varphi(r) = 1 + w(r)$. Zur Unendlichkeit der Variation vgl. S. 590 ebendort.

A KOORDINÁTA-MÉRÉS PONTOSSÁGÁNAK HATÁRAI AZ ELEKTRON DIRAC-FÉLE ELMÉLETÉBEN.

1. Az elektron állapotát a quantum-elmélet modern (HEISENBERG—DIRAC—SCHRÖDINGER-féle) alakjában az ú. n. hullámfüggvény, $\varphi(q)$ jellemzi; hol q , az elektron koordinátája, az argumentum, míg p , az impulzus nem szerepel.¹ $\varphi(q)$ helyett egy csak p -től, de q -tól nem, függő függvény is használható, $\psi(p)$. E két függvény a LAPLACE-féle transzformáció révén függ össze:

$$\psi(p) = \frac{1}{\sqrt{h}} \int \varphi(q) e^{\frac{2\pi i}{h} pq} dq, \quad \varphi(q) = \frac{1}{\sqrt{h}} \int \psi(p) e^{-\frac{2\pi i}{h} pq} dp.$$

A koordináta és az impulzus az elektron egy meghatározott állapotában általában nem bír határozott értékkel, hanem mindkettő több értéktartományban lehet pozitív valószínűségekkel. Annak a valószínűsége, hogy q a q' , $q' + dq'$ intervallumban fekszik, $|\varphi(q')|^2 \cdot dq'$; annak, hogy p a p' , $p' + dp'$ intervallumban fekszik, $|\psi(p')|^2 \cdot dp'$.³ (Mint hogy az összvalószínűségnek 1-nek kell lennie, $\int |\varphi(q)|^2 dq = \int |\psi(p)|^2 dp = 1$ egyenletnek kell fennállnia, — ez csakugyan minden hullámfüggvénynél igaz.)⁴

¹ PL. SCHRÖDINGER, «Abhandlungen zur Wellenmechanik», Leipzig 1927. Elsősorban a két első értekezés.

² V. ö. ¹ A következőkben h a PLANCK-féle állandó, m az elektron tömege, e a töltés, c a fénysebesség.

³ V. ö. ¹ Ezen képletek a quantummechanika ú. n. «statisztikai interpretáció»-jának legegyszerűbb alosztai. V. ö. BORN, Zschr. f. Phys., 37, 38, 1926; JORDAN, Zschr. f. Phys., 1926; DIRAC, Proc. Roy. Soc., 113, 1926; NEUMANN, Gött. Nachr., 1926.

⁴ V. ö. a fenti idézetek bármelyikével. Matematikai szempontból kimerítő PLANCHEREL, Circ. Math. Palermo, 1910-

Ezen képletek meghatározzák többek között q és p statisztikai középértékét és diszperzióját:⁵

$$\bar{q} = \int q |\varphi(q)|^2 dq, \quad D_q^2 = \overline{(q - \bar{q})^2} = \int (q - \bar{q})^2 |\varphi(q)|^2 dq,$$

$$\bar{p} = \int p |\psi(p)|^2 dp, \quad D_p^2 = \overline{(p - \bar{p})^2} = \int (p - \bar{p})^2 |\psi(p)|^2 dp.$$

q akkor veszi fel a $\bar{q}$ értéket igen pontosan (azaz: akkor fekszik túlnyomó valószínűséggel $\bar{q}$ -hoz igen közel), ha D_q kicsi, vagyis, ha $\varphi(q)$ csak a $\bar{q}$ hely közvetlen közelében vesz fel nagyobb értékeket; a megfelelő viszony áll fenn p , $\bar{p}$ és $\psi(p)$ között. De ha $\varphi(q)$ az említett viselkedést mutatja, akkor $\psi(p)$ ezt nem teheti⁶ és viszont — vagyis: kivételes esetekben q vagy p igen pontosan meg lehet határozva, de mind a kettő egyszerre soha.

HEISENBERG mutatott rá arra, hogy ez kísérleti szempontból annyit jelent, hogy nagy pontosságú koordináta- és impulzusmérés kizárja egymást.⁷ Úgy a q - és p -mérési kísérlete szemléletes diszkussziója, mint egyes tipikus hullámfüggvények D_q - és D_p -jének meghatározása segítségével arra a konklúzióra jutott, hogy D_q vagy D_p ugyan külön-külön tetszőlegesen kicsi lehet, de a $D_q D_p$ szorzat nagyságrendje mindig legalább is h . A hullámfüggvények általános matematikai vizsgálata által azonban pontos alsó határnak meg kell határozhatónak lennie, ez irányban KENNARD és ROBERTSON⁸ kimutatták, hogy $D_q D_p$ minimuma pontosan $\frac{h}{4\pi}$.

Mindezen megfontolások a relativitás elvét nem veszik figyelembe. Az elektron relativisztikus tömegváltozásának, valamint mechanikai és mágneses forgató nyomatékának, az ú. n. spinnek figyelembevétele lényegesen befolyásolja HEISENBERG-nek a mérésre vonatkozó megfontolásait. Több szerző ezen megfontolások revíziója alapján azon eredményre jutott, hogy a koor-

⁵ Ha r egy fizikai mennyiség, úgy $\bar{r}$ annak statisztikai középértéke és $\overline{(r - \bar{r})^2}$ a középértéktől való eltérés négyzetének középértéke, az ú. n. diszperzió négyzete.

⁶ Ha $\varphi(q)$ pl. csak $q=0$ közelében különbözik 0-tól, úgy $\psi(p)$ majdnem konstans és viszont.

⁷ HEISENBERG, Zschr. f. Phys., 1926; BOHR, Naturwiss., 1928.

⁸ KENNARD, Zschr. f. Phys., 1927; ROBERTSON, Phys. Rev., 1930.

dinátának $\frac{h}{mc}$ -nél² nagyobb pontossággal való mérése lehetetlen,⁹ vagyis, hogy D_q nagyságrendje mindig legalább is $\frac{h}{mc}$. Ezen állítás elméleti alátámasztására azonban szükség volna arra, hogy a nem-relativisztikus és spin-nélküli $\varphi(q)$ hullámfüggvényre vonatkozó KENNARD—ROBERTSON-féle számítások az elektron relativisztikus és a spint is figyelembevevő DIRAC-féle elméletének keretében¹⁰ megismételtessenek és azt az eredményt adják, hogy D_q nem lehet kisebb nagyságrendű $\frac{h}{mc}$ -nél.

A DIRAC-féle elméletben az elektron hullámfüggvénye q -n kívül még egy további ζ koordinátától is függ (ζ az ú. n. spin-koordináta és csak 4 diszkrét értéket vesz fel: $\zeta=1, 2, 3, 4$), vagyis $\varphi(q)$ helyébe $\varphi(q, \zeta)$ lép, hasonlóan $\psi(p)$ helyébe $\psi(p, \zeta)$ lép. A kettő közötti összefüggés a régi:

$$\psi(p, \zeta) = \frac{1}{\sqrt{h}} \int \varphi(q, \zeta) e^{\frac{2\pi i}{h} pq} dq,$$

$$\varphi(q, \zeta) = \frac{1}{\sqrt{h}} \int \psi(p, \zeta) e^{-\frac{2\pi i}{h} pq} dp,$$

és annak a valószínűsége, hogy q q' , $q' + dq'$ -ben, illetve hogy p p' , $p' + dp'$ -ben legyen, $\sum_1^4 |\varphi(q', \zeta)|^2 dq'$, illetve $\sum_1^4 |\psi(p', \zeta)|^2 dp'$ (a «normirozás» feltétele $\sum_1^4 \int |\varphi(q, \zeta)|^2 dq = \sum_1^4 \int |\psi(p, \zeta)|^2 dp = 1$).

A középérték és a diszperzió képletei:

$$\bar{q} = \sum_1^4 \int q |\varphi(q, \zeta)|^2 dq, \quad D_q^2 = \sum_1^4 \int (q - \bar{q})^2 |\varphi(q, \zeta)|^2 dq$$

$$\bar{p} = \sum_1^4 \int p |\psi(p, \zeta)|^2 dp, \quad D_p^2 = \sum_1^4 \int (p - \bar{p})^2 |\psi(p, \zeta)|^2 dp.$$

Amint a D_q -kifejezés szerkezete mutatja, semmi akadályja sincs annak, hogy D_q tetszőlegesen kicsi legyen: csupán úgy kell $\varphi(q, \zeta)$ -t megválasztani, hogy csak $\bar{q}$ közvetlen közelébe eső

⁹ PL. LANDAU és PEIERLS, Zschr. f. Phys., 1931.

¹⁰ DIRAC, Proc. Roy. Soc., 1928.

q mellett vegyen fel nagyobb értékeket. Ámde a DIRAC-féle elmélet általában elfogadott értelmezése szerint nem minden $\varphi(q, \zeta)$ hullámfüggvény felel meg az elektron fizikailag leírt állapotainak, hanem csak az ú. n. tiszta pozitív hullámfüggvények,¹¹ és könnyen belátható, hogy a fentemlített $\varphi(q, \zeta)$ -k általában nem tiszta pozitívak. Ennélfogva a voltaképpeni kérdés ez: minő értékeket vehetnek fel D_q , D_p és $D_q D_p$, ha $\varphi(q, \zeta)$ egy tetszőleges tiszta pozitív hullámfüggvény.

2. A DIRAC-féle elmélet az elektron viselkedését, hasonlóan a SCHRÖDINGER-féléhez, azáltal írja le, hogy megadja, mint változik az idők folyamán az elektron állapota (azaz a $\varphi(q, \zeta)$ vagy a $\psi(p, \zeta)$ hullámfüggvény), ha az elektromágneses tér adva van. Tehát a hullámfüggvény az időtől, t -től is függ: $\varphi(t, q, \zeta)$, illetve $\psi(t, p, \zeta)$; és változását egy t -ben elsőfokú differenciálegyenlet határozza meg (úgyhogy a kezdőállapot, vagyis a $t=0$ -hoz tartozó hullámfüggvény, minden egyebet meghatároz). Ha még figyelembe vesszük, hogy q voltaképpen a 3 kartézi koordinátát képviseli: q_1, q_2, q_3 , és megfelelően p -nek is 3 komponense van: p_1, p_2, p_3 (a LAPLACE-féle transzformáció kitevőjében álló pq helyébe $\sum_1^3 p_n q_n$ irandó), akkor DIRAC differenciálegyenlete ez:

$$\left[\frac{1}{c_i} \left(\frac{h}{2\pi i} \frac{\partial}{\partial t} - \frac{e}{c} \mathfrak{A}_0 \right) \alpha_0 + \sum_1^3 n \left(\frac{h}{2\pi i} \frac{\partial}{\partial p_n} - \frac{e}{c} \mathfrak{A}_n \right) \alpha_n + mc \right] \varphi = 0. \quad (1)$$

Itt $\alpha_0, \alpha_1, \alpha_2, \alpha_3$ 4 4-soros matrix, melyek az

$$\alpha_m \alpha_n + \alpha_n \alpha_m = \begin{cases} 2 \cdot \mathbf{1}, & \text{ha } m = n \\ \mathbf{0}, & \text{ha } m \neq n \end{cases} \quad (2)$$

«kommutációs» feltételeknek tesznek eleget,^{11 12} $\mathfrak{A}_0, \mathfrak{A}_1, \mathfrak{A}_2, \mathfrak{A}_3$

¹¹ V. ö. ¹⁰ továbbá NEUMANN, Zschr. f. Phys., 1928; SCHRÖDINGER, Berl. Ber., 1931.

¹² Ha α az $\{\alpha^{\zeta\zeta'}\}_{\zeta, \zeta' = 1, 2, 3, 4}$ matrix, úgy $\alpha\varphi(q, \zeta)$ alatt a

$$\varphi_1(q, \zeta) = \sum_1^4 \zeta' \alpha^{\zeta\zeta'} \varphi(q\zeta')$$

függvény értendő. Ugyanúgy $\alpha\psi(p, \zeta)$ a

$$\psi_1(p, \zeta) = \sum_1^4 \zeta' \alpha^{\zeta\zeta'} \psi(p, \zeta')$$

függvény.

pedig az elektromágneses teret meghatározó «4-es vektorpotenciál» komponensei.

Ha elektromágneses tér nincs jelen, úgy (1) a «szabad» elektron

$$\left[\frac{1}{c_i} \frac{h}{2\pi i} \frac{\partial}{\partial t} \alpha_0 + \sum_1^3 \frac{h}{2\pi i} \frac{\partial}{\partial q_n} \alpha_n + mci \right] \varphi = 0 \quad (1')$$

differenciálegyenletévé specializálódik. A «stacioner» megoldások, melyeknél φ t -től csak egy $e^{\frac{2\pi i}{h} Et}$ -alakú faktor révén függ, tehát melyeknél

$$\varphi = \varphi(t, q, \zeta) = e^{\frac{2\pi i}{h} Et} \varphi_0(q, \zeta), \quad \psi = \psi(t, p, \zeta) = e^{\frac{2\pi i}{h} Et} \psi_0(p, \zeta),$$

különösen fontosak, mert (2) minden megoldása belőlük lineárisan tevődik össze. Ezeknél (2), ha még a LAPLACE-féle transzformáció segítségével áttérünk ψ -re, így hangzik:

$$\left[\frac{E}{ci} \alpha_0 + \sum_1^3 p_n \alpha_n + mci \right] \psi_0 = 0. \quad (3)$$

Ezen egyenleten feltűnő, hogy differenciációkat nem tartalmaz, csupán $\psi_0(p, \zeta)$ -nak ugyanazon p -hez tartozó 4 értéke (melyek $\zeta = 1, 2, 3, 4$ -nek felelnek meg) között ír elő bizonyos lineáris relációkat. Ha a

$$\alpha(E, p) = \frac{E}{ci} \alpha_0 + \sum_1^3 p_n \alpha_n + mci \quad (4)$$

(4-edfokú) matrix rangja 4 (azaz determinánsa $\neq 0$), úgy (3) ezen E, p -kre $\psi_0 = 0$ -t ír elő valamennyi ζ -ra. Vagyis: adott E mellett csak ama p impulzusok bírhatnak pozitív valószínűséggel (melynek értéke $\sum_1^4 |\psi(t, p, \zeta)|^2 dp = \sum_1^4 |\psi_0(p, \zeta)|^2 dp$, hol $dp = dp_1 \cdot dp_2 \cdot dp_3$), melyeknél $\alpha(E, p)$ rangja < 4 . $\alpha(E, p)$ és a hozzá konjugált-transzponált matrix determinánsai komplex-konjugált számok, tehát a szorzat determinánsa az előbbi abszolút értékének négyzete. Mivel az α_n -ek konjugált-transzponált-jaikkal egyeznek,^{11 12} az $\alpha(E, p)$ -é

$$\beta(E, p) = -\frac{E}{c_i} \alpha_0 + \sum_1^3 p_n \alpha_n - mci = -\alpha(E, -p),$$

és a kettő szorzata, a kommutációs feltételek alapján

$$\left(-\frac{E^2}{c^2} + \sum_1^3 p_n^2 + m^2 c^2 \right) \mathbf{1},$$

ennek determinánsa $\left(-\frac{E^2}{c^2} + \sum_1^3 p_n^2 + m^2 c^2 \right)^4$. A keresett feltétel tehát ez:

$$-\frac{E^2}{c^2} + \sum_1^3 p_n^2 + m^2 c^2 = 0 \text{ vagy: } E = \pm c \sqrt{m^2 c^2 + \sum_1^3 p_n^2}. \quad (5)$$

Mint hogy E a quantum-elmélet általánosan elfogadott értelmezése szerint az energia,¹³ (5) nem egyéb, mint az energiának az impulzusoktól való függésének jól ismert relativisztikus alakja.

Ha (5) fennáll, $\alpha(E, p)$ rangja mindenesetre < 4 , könnyen kimutatható, hogy $= 2$.¹³ Ezen E, p -khez tehát még két lineárisan független ϕ_0 tartozik, ami ama fizikai tény aequivalense, hogy az adott impulzusú és energiájú elektronnak (a «monochromatikus» de BROGLIE-hullámnak) még kétféle spinje lehet. Ezzel szemben nem kielégítő az, hogy (5)-ben a négyzetgyök + előjele mellett annak — előjele is megengedett, vagyis, hogy $E > 0$ mellett $E < 0$ -hoz is tartoznak megoldások. Ezen $E < 0$ megoldások viselkedése egyébként akkor, amikor elektromágneses tér is jelen van (azaz erők hatnak az elektronra), a tapasztalatokkal összeegyeztethetetlen és súlyosan paradox, amint azt az (1) egyenlet diszkussziója mutatja.¹⁴ (Nevezetesen úgy viselkednek, mintha az elektron tömege m helyett $-m$ volna.) Ha az $E > 0$ megoldások lineárkombinációit tiszta pozitív és az $E < 0$ megoldások lineárkombinációit tiszta negatív hullámfüggvényeknek nevezzük, akkor (2), illetve (3) minden megoldása egyértelműen

¹³ V. Ö. NEUMANN, ¹⁰.

¹⁴ V. Ö. ¹¹ DIRAC, Proc. Roy. Soc., 1930, megkísérelte az elektron «negatív» állapotait a protonnal azonosítani, ezen felfogással szemben azonban súlyos aggályok merültek fel, például OPPENHEIMER, Phys. Review, 1930.

előállítható mint egy tiszta pozitív és egy tiszta negatív hullámfüggvény összege; és ha a $t = 0$ pillanatra szorítkozunk, úgy minden hullámfüggvényre érvényes ez az egyértelmű felbontás. A fent mondottak értelmében normális fizikai világképünkbe csak a tiszta pozitív hullámfüggvények illeszkednek be, ennél fogva csak ezeket ismerjük el az elektron fizikailag lehetséges állapotainak.

A tiszta pozitív állapotot (5) értelmében

$$E = +c \sqrt{m^2 c^2 + \sum_1^3 p_n^2}$$

jellemzi, vagyis a (3) egyenlet E ezen értékével, melyet $\left(\frac{ci}{E} \alpha_0\right)$ -val megszorozva és a kommutációs feltételek segítségével) így is írhatunk: $(1 - U) \phi_0 = 0$, vagyis $U \phi_0 = \phi_0$, hol

$$U = \frac{m c \alpha_0 - \sum_1^3 p_n i \alpha_0 \alpha_n}{\sqrt{m^2 c^2 + \sum_1^3 p_n^2}}. \quad (6)$$

A tiszta negatív állapotot ugyanezen feltétel jellemzi, de a négyzetgyök ellenkező előjelével, tehát U helyet $-U$ -val, vagyis $U \phi_0 = -\phi_0$. Ha a ϕ_0 -król lineárkombinációkra, az általános $\phi = \phi(p, \zeta)$ hullámfüggvényekre térünk át, úgy megkapjuk ezekre is a tiszta pozitív, illetve tiszta negatív jelleg feltételét:

$$U \phi = \phi, \quad \text{illetve} \quad = -\phi.$$

A LAPLACE-féle transzformáció segítségével áttérhetnénk $\varphi = \varphi(q, \zeta)$ -re, vagyis itt is kiszámíthatnók az U operátort, de ezen komplikált számításra nem lesz szükségünk.

Végül megjegyezzük, hogy a kommutációs feltételek segítségével (6)-ból $U^2 = 1$ következik, valamint az, hogy U saját konjugált transzponáltjával azonos (ez utóbbi azért áll U -ra, mert α_0 -ra és $i \alpha_0 \alpha_n$ -re áll, t. i. $i \alpha_0 \alpha_n$ konjugált transzponáltja $-i \alpha_n \alpha_0 = i \alpha_0 \alpha_n$ [$n = 1, 2, 3$]). Ennek alapján kitűnik, hogy a tetszőleges ϕ felbontása egy tiszta pozitív ϕ_+ -ra és egy tiszta negatív ϕ_- -ra ez:

$$\phi = \phi_+ + \phi_-, \quad \phi_+ = \frac{\phi + U \phi}{2}, \quad \phi_- = \frac{\phi - U \phi}{2}.$$

3. Mielőtt D_q -t és D_p -t vizsgáljuk, megállapítjuk a következőt: Minthogy a tér homogén, ugyanazon D_q , D_p értékek fordulnak elő az összes (tisztá pozitív) hullámfüggvényeknél, mint azoknál, melyeknél $\bar{q}=0$ (vagyis $\bar{q}_1=\bar{q}_2=\bar{q}_3=0$): hiszen a térnek $-\bar{q}_1, -\bar{q}_2, -\bar{q}_3$ -al való eltolása¹⁵ által mindig elérhetjük ezt. Minthogy ezen elmélet relativisztikusan invariáns,¹⁶ ugyanazon D_q , D_p értékek fordulnak elő az összes tisztá pozitív hullámfüggvénynél, mint azoknál, melyeknél $\bar{p}=0$ (vagyis $\bar{p}_1=\bar{p}_2=\bar{p}_3=0$), hiszen egy helyesen megválasztott LORENTZ-féle transzformáció¹⁶ által mindig elérhetjük ezt.¹⁷ Ha $\bar{q}_n=0$, $\bar{p}_n=0$ ($n=1, 2, 3$), úgy

$$D_{q_n}^2 = \sum_1^4 \int \int \int q_n^2 |\varphi(q, \zeta)|^2 dq_1 dq_2 dq_3,$$

$$D_{p_n}^2 = \sum_1^4 \int \int \int p_n^2 |\psi(p, \zeta)|^2 dp_1 dp_2 dp_3.$$

Minthogy a jobboldali kifejezések általában $\bar{q}_n^2 = (\overline{q_n - \bar{q}_n})^2 + \bar{q}_n^2 = D_{q_n}^2 + \bar{q}_n^2$ -al, illetve $\bar{p}_n^2 = (\overline{p_n - \bar{p}})^2 + \bar{p}_n^2 = D_{p_n}^2 + \bar{p}_n^2$ -al egyenlők, tehát mindenestre $\geq D_{q_n}^2$, illetve $\geq D_{p_n}^2$, ennél fogva a minimum-vizsgálatainknál el is tekinthetünk a $\bar{q}_n=0$, $\bar{p}_n=0$ feltételektől. Ha még figyelembe vesszük, hogy a LAPLACE-féle transzformáció, mely φ -t ψ -vé, $q_n\varphi$ -t $\frac{h}{2\pi i} \frac{\partial}{\partial p_n} \psi$ -vé alakítja át, úgyhogy ezeknek abszolútérték-négyzetintegráljait meg nem változtatja, úgy a következőt látjuk: D_{q_n} , D_{p_n} megfontolásainkban pótolhatók a $\left(\frac{h}{2\pi}\right) E_n$, F_n mennyiségekkel, hol

¹⁵ Ezen eltolás $\varphi(q_1, q_2, q_3, \zeta)$ -nak $\varphi(q_1 - \bar{q}_1, q_2 - \bar{q}_2, q_3 - \bar{q}_3, \zeta)$ -val és $\psi(p_1, p_2, p_3, \zeta)$ -nak $e^{\frac{2\pi i}{h}(\bar{q}_1 p_1 + \bar{q}_2 p_2 + \bar{q}_3 p_3)} \psi(p_1, p_2, p_3, \zeta)$ -val való pótlását jelenti.

¹⁶ V. ö. 10

¹⁷ Ezután ¹⁵ nyomán elérhetjük $\bar{q}=0$ -t, ezen művelet

$$\bar{p}_n = \sum_1^4 \int \int \int p_n |\psi(p_1, p_2, p_3, \zeta)|^2 dp_1 dp_2 dp_3$$

értékét nyilván nem érinti, tehát $\bar{p}=0$ érvényben marad.

$$E_n^2 = \sum_{\zeta} \int \int \int \left| \frac{\partial}{\partial p_n} \phi(p, \zeta) \right|^2 dp_1 dp_2 dp_3, \\ F_n^2 = \sum_{\zeta} \int \int \int |p_n \phi(p, \zeta)|^2 dp_1 dp_2 dp_3. \quad (7)$$

Mindezen egyenletekben, éppúgy, mint az U operátor kifejezésében egyedül $\phi(p, \zeta)$ szerepel ($\varphi(q, \zeta)$ nem!). Természetesen a

$$\sum_{\zeta} \int \int \int |\phi(p, \zeta)|^2 dp_1 dp_2 dp_3 = 1 \quad (7')$$

mellékfeltétel (mely $\sum_{\zeta} \int \int \int |\varphi(q, \zeta)|^2 dq_1 dq_2 dq_3 = 1$ -et a fentemlített oknál fogva maga után vonja) ismét elő van írva.

A következő rövidítéseket vezetjük be:

$$(\phi, \omega) = \sum_{\zeta} \int \int \int \phi(p, \zeta) \overline{\omega(p, \zeta)} dp_1 dp_2 dp_3,^{18}$$

$$\|\phi\| = + \sqrt{(\phi, \phi)} = + \sqrt{\sum_{\zeta} \int \int \int |\phi(p, \zeta)|^2 dp_1 dp_2 dp_3}.$$

Ezáltal (7), (7') így alakul át:

$$E_n = \left\| \frac{\partial}{\partial p_n} \phi \right\|, \quad F_n = \|p_n \phi\|, \quad \|\phi\| = 1. \quad (8)$$

Ha α_0 -t $-\alpha_0$ -val pótoljuk, a definiáló (2) egyenlet érvényben marad, tehát a DIRAC-féle egyenlet általános tulajdonságai nem változnak meg. Mivel ekkor $U = -U$ -vá alakul át, a (8)-beli mennyiségek $U\psi = \psi$ megoldásainál ugyanúgy viselkednek, mint $U\psi = -\psi$ megoldásainál. Vagyis a tiszta pozitív ϕ -k helyett a tiszta negatív ϕ -ket is vizsgálhatjuk.

¹⁸ Itt — és ezután is a komplex-konjugált számot jelöli, a középértéket jelölő —-el nem tévesztendő össze!

A következőkben a ψ -hez tartozó E_n, F_n -t a $\psi_1 = \frac{\psi_+}{\|\psi_+\|}$ és $\psi_2 = \frac{\psi_-}{\|\psi_-\|}$ -hoz¹⁹ tartozókkal fogjuk összehasonlítani; ψ_1, ψ_2

ama tiszta pozitív, illetve negatív hullámfüggvények, melyekből ψ lineárisan összetevődik. A későbbiekben még használni fogjuk a következő megjegyzést:

$$\begin{aligned}\|\psi_+\|^2 &= \left(\frac{1+U}{2} \psi, \frac{1+U}{2} \psi \right) = \left(\left(\frac{1+U}{2} \right)^2 \psi, \psi \right) = \\ &= \left(\frac{1+2U+U^2}{4} \psi, \psi \right) = \left(\frac{1+U}{2} \psi, \psi \right),\end{aligned}$$

$$\|\psi_-\|^2 = \dots = \left(\frac{1-U}{2} \psi, \psi \right),$$

$$\|\psi_+\|^2 + \|\psi_-\|^2 = (\psi, \psi) = 1,$$

tehát $\|\psi_+\|$ és $\|\psi_-\|$ közül az egyik $\geq \frac{1}{\sqrt{2}}$.

4. Amint kimutattuk, $\|\psi_+\|^2 + \|\psi_-\|^2 = \|\psi\|^2$ valamennyi ψ -re, tehát $\|\psi_+\| \leq \|\psi\|$. Mivel (6), (8) alapján nyilván $p_n \psi_\pm = (p_n \psi)_\pm$, $\|p_n \psi_+\| \leq \|p_n \psi\| = F_n$, tehát

$$F'_n = \|p_n \psi_1\| = \frac{\|p_n \psi_+\|}{\|\psi_+\|} \leq \frac{F_n}{\|\psi_+\|}. \quad (9')$$

Ugyanígy

$$F''_n = \|p_n \psi_2\| = \dots \leq \frac{F_n}{\|\psi_-\|}.$$

Voltaképpen feladatunk tehát $E'_n = \left\| \frac{\partial}{\partial p_n} \psi_1 \right\|$ és $E''_n = \left\| \frac{\partial}{\partial p_n} \psi_2 \right\|$ felső határait megtalálni.

¹⁹ ψ_1, ψ_2 a ψ_+, ψ_- -al arányos, de normirozott függvények: e nélkül nem volnának hullámfüggvénynek használhatók.

Mármost

$$\left. \begin{aligned}
 E'_n &= \left\| \frac{\partial}{\partial p_n} \phi_+ \right\| = \frac{\left\| \frac{\partial}{\partial p_n} \phi_+ \right\|}{\left\| \phi_+ \right\|}, \quad E''_n = \left\| \frac{\partial}{\partial p_n} \phi_- \right\| = \frac{\left\| \frac{\partial}{\partial p_n} \phi_- \right\|}{\left\| \phi_- \right\|}, \\
 \left\| \frac{\partial}{\partial p_n} \phi_+ \right\|^2 &= \left(\frac{\partial}{\partial p_n} \phi_+, \frac{\partial}{\partial p_n} \phi_+ \right) = \left(-\frac{\partial^2}{\partial p_n^2} \phi_+, \phi_+ \right) = \\
 &= \left(\frac{\partial^2}{\partial p_n^2} \frac{1+U}{2} \phi, \frac{1+U}{2} \phi \right) = \\
 &= \left(-\frac{1+U}{2} \frac{\partial^2}{\partial p_n^2} \frac{1+U}{2} \phi, \phi \right)^{20} \\
 \left\| \frac{\partial}{\partial p_n} \phi_- \right\|^2 &= \dots = \left(-\frac{1-U}{2} \frac{\partial^2}{\partial p_n^2} \frac{1-U}{2} \phi, \phi \right), \\
 \left\| \frac{\partial}{\partial p_n} \phi_+ \right\|^2 + \left\| \frac{\partial}{\partial p_n} \phi_- \right\|^2 &= \\
 &= \left(-\left[\frac{1+U}{2} \frac{\partial^2}{\partial p_n^2} \frac{1+U}{2} + \frac{1-U}{2} \frac{\partial^2}{\partial p_n^2} \frac{1-U}{2} \right] \phi, \phi \right) = \\
 &= \left(-\frac{1}{2} \left(\frac{\partial^2}{\partial p_n^2} + U \frac{\partial^2}{\partial p_n^2} U \right) \phi, \phi \right) = \\
 &= \left(-\frac{\partial^2}{\partial p_n^2} \phi, \phi \right) + \left(\frac{1}{2} \left(\frac{\partial^2}{\partial p_n^2} - U \frac{\partial^2}{\partial p_n^2} U \right) \phi, \phi \right).
 \end{aligned} \right\} (10)$$

A jobboldal első tagja nyilván $= \left(\frac{\partial}{\partial p_n} \phi, \frac{\partial}{\partial p_n} \phi \right) = E_n^2$, tehát csupán a másodikkal kell törődnünk. A benne szereplő

$$\frac{\partial^2}{\partial p_n^2} - U \frac{\partial^2}{\partial p_n^2} U = \frac{\partial}{\partial p_n} U \cdot U \frac{\partial}{\partial p_n} - U \frac{\partial}{\partial p_n} \cdot \frac{\partial}{\partial p_n} U = V_n$$

operátort a $\frac{\partial}{\partial p_n} U$ és $U \frac{\partial}{\partial p_n}$ operátorok kommutátorának tekintetjük, vagy $W_n = \frac{\partial}{\partial p_n} U - U \frac{\partial}{\partial p_n}$ és $U \frac{\partial}{\partial p_n}$ kommutátorának. Minthogy U nem tartalmaz differenciációkat

$$W_n = \frac{\partial}{\partial p_n} U - U \frac{\partial}{\partial p_n} = \frac{\partial}{\partial p_n} (U);^{21}$$

²⁰ E helyen használhatjuk ama közismert tényt, hogy $\frac{\partial}{\partial p_n}$ ú. n. ferde-Hermitei szimmetriájú operátor, míg U Hermitei szimmetriájú.

²¹ Ezen fontos kommutációs képlet eredete: BORN, HEISENBERG és JORDAN, Zschr. f. Phys., 1926, és SCHRÖDINGER, ¹, harmadik értekezés.

tehát W_n sem tartalmaz differenciációkat és így U -val kommutál. W_n és $U \frac{\partial}{\partial p_n}$ kommutátora tehát U -szor W_n és $\frac{\partial}{\partial p_n}$ kommutátora, mely utóbbi

$$-\frac{\partial}{\partial p_n} (W_n)^{21} = -\frac{\partial^2}{\partial p_n^2} (U).$$

A szóban forgó operátor ennélfogva

$$V_n = -U \frac{\partial^2}{\partial p_n^2} (U) = \frac{1}{2} \left[\left(\frac{\partial}{\partial p_n} (U) \right)^2 - \frac{\partial^2}{\partial p_n^2} (U^2) \right].$$

Legyen pl. $n=3$, (6) értelmében

$$\begin{aligned} & \frac{\partial}{\partial p_3} (U) = - \\ & = - \frac{p_3}{\sqrt{m^2 c^2 + \sum_1^3 p_n^2}} \left(m c \alpha_0 - \sum_1^3 p_n i \alpha_0 \alpha_n \right) - \frac{i \alpha_0 \alpha_3}{\sqrt{m^2 c^2 + \sum_1^3 p_n^2}} = \\ & = \frac{m c p_3 \alpha_0 + p_1 p_3 i \alpha_0 \alpha_1 + p_2 p_3 i \alpha_0 \alpha_2 - (m^2 c^2 + p_1^2 + p_2^2) i \alpha_0 \alpha_3}{\sqrt{m^2 c^2 + \sum_1^3 p_n^2}} \end{aligned}$$

ha ezt (2) figyelembevételével négyzetre emeljük,²² ez adódik:

$$\begin{aligned} & \left(\frac{\partial}{\partial p_3} (U) \right)^2 = \frac{m^2 c^2 p_3^2 + p_1^2 p_3^2 + p_2^2 p_3^2 + (m^2 c^2 + p_1^2 + p_2^2)^2}{(m^2 c^2 + p_1^2 + p_2^2 + p_3^2)^3} = \\ & = \frac{m^2 c^2 + p_1^2 + p_2^2}{(m^2 c^2 + p_1^2 + p_2^2 + p_3^2)^2}. \end{aligned}$$

Másfelől tudjuk, hogy $U^2=1$, tehát $\frac{\partial^2}{\partial p_3^2} (U^2) = 0$. Tehát

$$\begin{aligned} V_3 &= \frac{1}{2} \frac{m^2 c^2 + p_1^2 + p_2^2}{(m^2 c^2 + p_1^2 + p_2^2 + p_3^2)^2} \\ (V_3 \psi, \psi) &= \frac{1}{2} \sum_1^4 \iiint \frac{m^2 c^2 + p_1^2 + p_2^2}{(m^2 c^2 + p_1^2 + p_2^2 + p_3^2)^2} |\psi|^2 dp_1 dp_2 dp_3 \leq \\ &\leq \frac{1}{2} \sum_1^4 \iiint \frac{|\psi|^2 dp_1 dp_2 dp_3}{m^2 c^2 + p_1^2 + p_2^2 + p_3^2}, \end{aligned}$$

²² $\alpha_0, i\alpha_0\alpha_1, i\alpha_0\alpha_2, i\alpha_0\alpha_3$ kommutációs tulajdonságai egyeznek $\alpha_0, \alpha_1, \alpha_2, \alpha_3$ -éival. V. ö. ¹⁰

és így (tetszőleges $n = 1, 2, 3$ -ra):

$$\left\| \frac{\partial}{\partial p_n} \phi_+ \right\|^2 + \left\| \frac{\partial}{\partial p_n} \phi_- \right\|^2 \leq E_n^2 + \frac{1}{4} \sum_1^4 \iiint \frac{|\phi|^2 dp_1 dp_2 dp_3}{m^2 c^2 + \sum_1^3 p_n^2}. \quad (11)$$

5. Ha $\phi'_n = \frac{\phi}{p_n}$ szintén véges abszolútérték-négyzetintegrállal bír (vagyis: ha $p_n \rightarrow 0$ -ból $\phi \rightarrow 0$ következik, és pedíg elég erősen),²³ úgy (11) jobboldalának második tagját a következőképpen alakíthatjuk át:

$$\begin{aligned} \frac{1}{4} \sum_1^4 \iiint \frac{|\phi|^2 dp_1 dp_2 dp_3}{m^2 c^2 + \sum_1^3 p_n^2} &= \frac{1}{4} \sum_1^4 \iiint \frac{p_n^2 |\phi'_n|^2 dp_1 dp_2 dp_3}{m^2 c^2 + \sum_1^3 p_n^2} \leq \\ &\leq \frac{1}{4} \sum_1^4 \iiint |\phi'_n|^2 dp_1 dp_2 dp_3 = \frac{1}{4} \|\phi'_n\|^2. \end{aligned}$$

Másfelől

$$\begin{aligned} 2\Re \left(\frac{\partial}{\partial p_n} \phi, \phi'_n \right) &= -2\Re \left(\phi, \frac{\partial}{\partial p_n} \phi'_n \right) = -2\Re \left(p_n \phi'_n, \frac{\partial}{\partial p_n} \phi'_n \right) = \\ &= - \left(p_n \phi'_n, \frac{\partial}{\partial p_n} \phi'_n \right) - \left(\frac{\partial}{\partial p_n} \phi'_n, p_n \phi'_n \right) = \\ &= \left(\frac{\partial}{\partial p_n} (p_n \phi'_n), \phi'_n \right) - \left(p_n \frac{\partial}{\partial p_n} \phi'_n, \phi'_n \right) = (\phi'_n, \phi'_n) = \|\phi'_n\|^2,^{20} \\ 2\Re \left(\frac{\partial}{\partial p_n} \phi, \phi'_n \right) &\leq 2 \left| \left(\frac{\partial}{\partial p_n} \phi, \phi'_n \right) \right| \leq 2 \left\| \frac{\partial}{\partial p_n} \phi \right\| \cdot \|\phi'_n\|,^{24} \end{aligned}$$

tehát

$$\|\phi'_n\|^2 \leq 2 \left\| \frac{\partial}{\partial p_n} \phi \right\| \cdot \|\phi'_n\|, \quad \|\phi'_n\| \leq 2 \left\| \frac{\partial}{\partial p_n} \phi \right\| = 2E_n.^{25}$$

²³ Hogy $\int \left| \frac{\psi}{p_n} \right|^2 dp_n$ véges legyen, elegendő, ha ψ aszimptotikusan p_n^a , hol $a > \frac{1}{2}$.

²⁴ Ez az ú. n. SCHWARZ-féle egyenlőtlenségből következik, mely szerint $|\langle \psi, \omega \rangle| \leq \|\psi\| \cdot \|\omega\|$

²⁵ Vagyis: $\|\psi'_n\|$ lehet ugyan végtelen, de ha véges, úgy $\leq \sqrt{2E_n}$.

Tehát a szóbanforgó második tag $\leq \frac{1}{4}(2E_n)^2 = E_n^2$, és (11)-ből ez lesz:

$$\left\| \frac{\partial}{\partial p_n} \phi_+ \right\|^2 + \left\| \frac{\partial}{\partial p_n} \phi_- \right\|^2 \leq 2E_n^2. \quad (12)$$

Ennélfogva $\left\| \frac{\partial}{\partial p_n} \phi_+ \right\|$ és $\left\| \frac{\partial}{\partial p_n} \phi_- \right\|$ külön-külön $\leq \sqrt{2} E_n$ és (10) értelmében

$$E'_n \leq \frac{\sqrt{2}}{\|\phi_+\|} E_n, \quad E''_n \leq \frac{\sqrt{2}}{\|\phi_-\|} E_n. \quad (13)$$

Amint 3. végén megjegyeztük, $\|\phi_+\|$ vagy $\|\phi_-\| \geq \frac{1}{\sqrt{2}}$, tehát E'_n vagy $E''_n \leq 2E_n$. Ugyanekkor (9'), (9'') értelmében F'_n , illetve $(F''_n \leq \sqrt{2} F_n$.

Ha tehát $p_n \rightarrow 0$ -ból $\phi \rightarrow 0$ következik, és pedig elég erősen,²³ úgy vagy a tiszta pozitív ϕ_+ -hoz tartozó E'_n , F'_n , vagy a tiszta negatív ϕ_- -hoz tartozó E''_n , F''_n nem lényegesen nagyobb, mint a ϕ -hez tartozó E_n , F_n . Ennélfogva ilyenkor a koordinata diszperzióját mérő E'_n , E''_n lényegében ugyanoly kicsivé tehető, mint az az eredeti E_n és a HEISENBERG—KENNARD—ROBERTSON-féle relációk maradnak érvényben tiszta pozitív, illetve negatív ϕ kre is. Egy a fenti feltételt kielégítő ϕ , melynél ezen relációk által megengedett legkedvezőbb eset ($E_n F_n \sim 1$) megközelítettetik, pl. ez:

$$\phi = C p_1 p_2 p_3 e^{-a(p_1^2 + p_2^2 + p_3^2)}.$$

$a > 0$, $C > 0$, C úgy választandó, hogy $\|\phi\| = 1$ legyen.) Könnyen belátható, hogy

$$\varphi = C' q_1 q_2 q_3 e^{-\frac{\pi^2}{h^2 a} (q_1^2 + q_2^2 + q_3^2)}$$

és $E_n = C'' a^{\frac{1}{2}}$, $F_n = C''' a^{-\frac{1}{2}}$ (C' , C'' , C''' a -t és h -t nem tartalmazó numerikus állandók), tehát $E_n F_n = C'' C'''$.

Másfelől a fenti feltételt ki nem elégítő ϕ -kre (11) jobb oldalának második tagja lényegesen nagyobb lehet, mint E_n^2 . Így a

$$\phi = D e^{-a p_1^2 - p_2^2 - p_3^2}$$

($a > 0$, $D > 0$, D úgy választandó, hogy $\|\phi\| = 1$ legyen) hullámfüggvénynél E_1 ismét $= D' a^{-\frac{1}{2}}$, tehát $E_1^2 = D'^2 a^{-1}$, ezzel szem-

ben $\frac{1}{4} \sum_{\zeta} \iiint \frac{|\psi|^2 dp_1 dp_2 dp_3}{m^2 c^2 + p_1^2 + p_2^2 + p_3^2}$ $a \rightarrow \infty$ esetén szintén 0-hoz konvergál ugyan, de, mint könnyen belátható, oly módon, mint $Da^{-\frac{1}{2}}$, tehát E_1^2 -nél lassabban.

Ennélfogva befejezésül ezen kifejezésnek E_n -hez való viszonyát az általános esetben kell tisztáznunk, midőn $p_n \rightarrow 0$ nem vonja szükségszerűen maga után $\phi \rightarrow 0$ -t.

6. Be fogjuk bizonyítani, hogy az A numerikus állandó alkalmas értéke mellett

$$\frac{1}{4} \sum_{\zeta} \iiint \frac{|\psi|^2 dp_1 dp_2 dp_3}{m^2 c^2 + p_1^2 + p_2^2 + p_3^2} \leq A \text{ Min}(E_1, E_2, E_3). \quad (14)$$

$n=1, 2, 3$ szimmetriájára való tekintettel elég $\leq AE_1$ -et bebizonyítani, vagyis, ha E_1 (7)-beli értékét behelyettesítjük és a $\|\phi\|=1$ korlátozást elejtjük:

$$\sum_{\zeta} \iiint \frac{|\psi|^2 dp_1 dp_2 dp_3}{m^2 c^2 + p_1^2 + p_2^2 + p_3^2} \leq 4A \left\| \frac{\partial}{\partial p_1} \phi \right\| \cdot \|\phi\|. \quad (14')$$

A jobboldal $u \left\| \frac{\partial}{\partial p_1} \phi \right\|^2 + v \|\phi\|^2$ minimuma, ha u, v az összes $u > 0, v > 0, uv = (4A)^2$ -ot kielégítő értékpárokat jelenti, a baloldalt növeljük, ha a nevezőt $m^2 c^2 + p_1^2$ -el pótoljuk. Ezenfelül

$$\left\| \frac{\partial}{\partial p_1} \phi \right\|^2 = \left(\frac{\partial}{\partial p_1} \phi, \frac{\partial}{\partial p_1} \phi \right) = \left(-\frac{\partial^2}{\partial p_1^2} \phi, \phi \right),^{20} \|\phi\| = (\phi, \phi)$$

tehát (14') helyett ezt is írhatjuk:

$$\begin{aligned} \left(\frac{1}{m^2 c^2 + p_1^2} \phi, \phi \right) &\leq \left(-u \frac{\partial^2}{\partial p_1^2} \phi + v \phi, \phi \right) \\ \left(\left(-u \frac{\partial^2}{\partial p_1^2} - \frac{1}{m^2 c^2 + p_1^2} \right) \phi, \phi \right) &\geq -v(\phi, \phi). \end{aligned} \quad (14'')$$

Mivel a baloldalon álló operátor a p_2, p_3 és ζ változókat nem érinti, ezektől eltekinthetünk. Ha p_1 helyébe $m c x$ -et írunk és (14'')-et megszorozzuk $\frac{m^2 c^2}{u}$ -val, úgy ez az egyenlőtlenség keletkezik:

$$\left(\left(-\frac{\partial^2}{\partial x^2} - \frac{1}{u(1+x^2)} \right) \phi, \phi \right) \geq -\frac{m^2 c^2 v}{u} (\phi, \phi).$$

Vagyis, ha $\frac{1}{u}$ helyébe a -t írunk és $v = \frac{(4A)^2}{u} = (4A)^2 a$ -t behelyettesítjük:

$$-\int_{-\infty}^{+\infty} \frac{\partial^2}{\partial x^2} \phi(x) \cdot \overline{\phi(x)} dx - a \int_{-\infty}^{+\infty} \frac{|\phi'(x)|^2}{1+x^2} dx \geq -Ba^2 \int_{-\infty}^{+\infty} |\phi(x)|^2 dx, \quad (15)$$

hol $B = (4mcA)^2$. Nyilván elég, ha $\int_{-\infty}^{+\infty} |\phi(x)|^2 dx = 1$ -re szorítkozunk, tehát ez esetben bebizonyítandó, hogy a baloldal minimuma $\geq -Ba^2$. Minthogy ϕ -jeinket approximálhatjuk oly ϕ -kkel, melyek egy tetszőleges, de véges intervallumon kívül eltűnnek, a szóbanforgó minimum $= \lim_{A \rightarrow \infty} M_A$, hol M_A (15) baloldalának minimuma, a

$$\phi(x) = 0, \quad \text{ha} \quad |x| \geq A$$

mellékfeltételek mellett. Ha N_A

$$-\int_{-A}^{+A} \frac{\partial^2}{\partial x^2} \phi(x) \overline{\phi(x)} dx - a \int_{-A}^{+A} \frac{|\phi'(x)|^2}{1+x^2} dx \quad (16)$$

minimuma a

$$\int_{-A}^{+A} |\phi(x)|^2 dx = 1, \quad \phi(\pm A) = 0 \quad (16')$$

mellékfeltételek mellett ($|x| > A$ -ra $\phi(x)$ definiálatlan lehet), úgy nyilván $M_A \geq N_A$. Tehát csak ezt kell bebizonyítanunk:

$$\lim_{A \rightarrow \infty} N_A \geq -Ba^2$$

(alkalmasan választott B -re, amely A -t meghatározza: $A = \frac{\sqrt{B}}{4mc}$).

(16), (16') egy reguláris variációs probléma,²⁶ melynél (16) minimumát, N_A -t, felveszi egy és (lényegében) csak egy ϕ -nél. Ezen $\phi(x) = \phi_A(x)$ kielégíti a (16) határfeltételeket, az EULER-féle differenciálegyenletet

$$-\frac{\partial^2}{\partial x^2} \phi(x) - a \frac{\phi(x)}{1+x^2} = \lambda \phi(x), \quad (17)$$

²⁶ PL. COURANT-HILBERT, «Methoden der mathematischen Physik, I.» Berlin, 1924. különösen a VI. fejezet.

és az $|x| < A$ intervallumban sehol el nem tűnik. Ennélfogva előírhatjuk, hogy $\phi_A(0) > 0$ legyen, és mivel $\phi_A(-x)$ szintén megoldása a minimum-feladatnak, $\phi_A(-x) = \pm \phi_A(x)$, és $x = 0$ behelyettesítése után $\phi_A(-x) = \phi_A(x)$. Az $\int_{-A}^{+A} \dots \overline{\phi(x)} dx$ műveletnek (17)-re való alkalmazása mutatja, hogy $\lambda = N_A$.

$\phi(x) = 1 - \frac{x^2}{A^2}$ behelyettesítése mutatja, hogy (16) (elég nagy A mellett) negatív értékeket is felvesz, másfelől (16) első tagja

$$\int_{-A}^{+A} \left| \frac{\partial}{\partial x} \phi(x) \right|^2 dx \geq 0,$$

a második pedig

$$> -a \int_{-A}^{+A} |\phi(x)|^2 dx = -a,$$

ennélfogva

$$\lambda = N_A < 0, > -a$$

(17) értelmében $\frac{\partial^2}{\partial x^2} \phi_A(x) = \left(-\frac{a}{1+x^2} - \lambda \right) \phi_A(x)$, és minthogy $\phi_A(x) > 0$, $\frac{\partial^2}{\partial x^2} \phi_A(x) \geq 0$, a szerint, hogy $x \leq x_0$, hol x_0 -t $-\frac{a}{\frac{\partial}{\partial x} \phi_A(x)} - \lambda = 0$, $a = \sqrt{\frac{a}{-\lambda} - 1} > 0$ határozza meg. Tehát $\frac{\partial}{\partial x} \phi_A(x)$ 0-tól x_0 -ig fogy, x_0 -tól kezdve nő. Minthogy 0-nál 0, A -nál < 0 a $0 < x \leq A$ intervallumban mindenesetre < 0 , minimuma x_0 -nál vagy A -nál fekszik. $\phi_A(x)$ tehát 0-tól A -ig fogy.

Mindezek alapján 0 és A között, ha $C = \phi_A(0)$,

$$\frac{\partial^2}{\partial x^2} \phi_A(x) \geq -\frac{Ca}{1+x^2},$$

$$\frac{\partial}{\partial x} \phi_A(x) \geq -\int_0^x \frac{Ca}{1+x^2} dx \geq -\int_0^\infty \frac{Ca}{1+x^2} dx = -\frac{\pi}{2} Ca,$$

$$\phi_A(x) \geq C - \int_0^x \frac{\pi}{2} Ca dx = C \left(1 - \frac{\pi}{2} ax \right).$$

Ennélfogva :

$$\frac{\int_0^{\frac{1}{\pi\alpha}} \frac{|\phi_A(x)|^2}{1+x^2} dx}{\int_0^{\frac{1}{\pi\alpha}} |\phi_A(x)|^2 dx} \leq \frac{\int_0^{\frac{1}{\pi\alpha}} C^2 \frac{dx}{1+x^2}}{\int_0^{\frac{1}{\pi\alpha}} \left(\frac{C}{2}\right)^2 dx} \leq \frac{C^2 \frac{\pi}{2}}{\left(\frac{C}{2}\right)^2 \frac{1}{\pi\alpha}} = 2\pi^2\alpha,$$

és

$$\frac{\int_{\frac{1}{\pi\alpha}}^A \frac{|\phi_A(x)|^2}{1+x^2} dx}{\int_{\frac{1}{\pi\alpha}}^A |\phi_A(x)|^2 dx} \leq \frac{1}{1 + \left(\frac{1}{\pi\alpha}\right)^2} = \frac{\pi^2\alpha^2}{\pi^2\alpha^2 + 1} < \pi^2\alpha^2 \leq 2\pi^2\alpha,$$

ha $\alpha \leq 2$, tehát

$$\frac{\int_0^A \frac{|\phi_A(x)|^2}{1+x^2} dx}{\int_0^A |\phi_A(x)|^2 dx} \leq 2\pi^2\alpha.$$

Minthogy a baloldal mindenestre < 1 , ez az egyenlőtlenség $\alpha > 2$ mellett is érvényes.

Ha $\int_0^A -t \int_{-A}^{+A}$ -val pótoljuk, a számláló és a nevező megduplázódik. A nevező ez esetben 1, tehát

$$\int_{-A}^{+A} \frac{|\phi_A(x)|^2}{1+x^2} dx \leq 2\pi^2\alpha.$$

Minthogy (16) első tagja parciális integráció útján $\int_{-A}^{+A} \left| \frac{\partial}{\partial x} \phi_A(x) \right|^2 dx$ -é alakítható át, tehát ≥ 0 és a második a fentiek szerint $\geq -2\pi^2\alpha^2$, bebizonyítottuk, hogy $N_A \geq -2\pi^2\alpha^2$. Tehát $B = 2\pi^2$, $A = \frac{\pi}{\sqrt{8mc}}$ megfelel kívánalmainknak.

7. (10), (11), (14) és A fenti értéke alapján

$$E'_n \leq \frac{\sqrt{E_n^2 + \frac{\pi}{\sqrt{8}mc} \text{Min}(E_1, E_2, E_3)}}{\|\phi_+\|}, \quad (18)$$

$$E''_n \leq \frac{\sqrt{E_n^2 + \frac{\pi}{\sqrt{8}mc} \text{Min}(E_1, E_2, E_3)}}{\|\phi_-\|}.$$

Tehát a szerint, hogy $\|\phi_+\|$ vagy $\|\phi_-\| \geq \frac{1}{\sqrt{2}}$ (v. ö. 3. végével)

E'_n , illetve $E''_n \leq \sqrt{2E_n^2 + \frac{\pi}{\sqrt{8}mc} \text{Min}(E_1, E_2, E_3)}$ és (v. ö. (9'), (9'')) F'_n , illetve $F''_n \leq \sqrt{2} F_n$. Tehát az általános esetben is E'_n 0-hoz tart E_n -el, de gyöngébben, mint E_n , illetve F''_n aszimptotikus viselkedése $\sqrt{E_n}$. Megjegyzendő, hogy a második tag, amely ezt okozza, akkor kezd az elsőn túlnőni, ha $\frac{\pi}{\sqrt{8}mc} E_n \geq E_n^2$, vagyis $E_n \geq \frac{\pi}{\sqrt{8}mc}$. Vagyis, ha a koordináta-mérés pontossága

$$\bar{q}_n = \frac{h}{2\pi} E_n \geq \frac{1}{\sqrt{32}} \frac{h}{mc}.$$

Mint tudjuk, a koordináta mérhetősége ellen kifejezett aggályok csakugyan mindig arra az esetre vonatkoztak, midőn a kívánt pontosság $\frac{h}{mc}$ vagy még több.

ÜBER DIE GRENZEN DER KOORDINATENMESSUNGS-GENAUIGKEIT IN DER DIRACSCHEN THEORIE DES ELEKTRONS

Die bekannten Untersuchungen Heisenbergs haben gezeigt, dass man auch ohne Verwendung der modernen Quantenmechanik, unter alleiniger Verwendung der qualitativen Grundeigenschaften der typischen Quantenprozesse (Compton-Effekt, Lichtemission in $h\nu$ -Quanten, u. Ä.), die «Unbestimmtheitsrelation» $\Delta p \cdot \Delta q > h$ herleiten kann. Dabei sind Δp , Δq die mittleren Fehler von zwei gleichzeitigen Impuls- und Koordinatenmessungen. Die Quantenmechanik bestätigt dies, und präzisiert es zu $\Delta p \cdot \Delta q \geq \frac{h}{4\pi}$. Alle diese Rechnungen sind unrelativistisch, und es wurde mehrfach der Verdacht ausgesprochen, dass eine relativistische Rechnung sogar eine absolute untere Grenze für Δq ergeben würde. Insbesondere wurde eine Bestätigung dieser Vermutung durch die sog. rein positiven Wellenfunktionen der DIRACschen Elektronentheorie erhofft. In der vorliegenden Arbeit wird gezeigt, dass dem nicht so ist, sondern dass bei diesen im Wesentlichen dieselbe «Unbestimmtheitsrelation» gilt, wie bei HEISENBERG.

ERRATA for

On an Algebraic Generalization of the Quantum Mechanical Formalism

p. 414, line 8 should read:

then⁷ $ab \in (N_{\frac{1}{2}} + N_1)$; if $a \in N_{\frac{1}{2}}$ and $b \in N_1$, then⁷ $ab \in (N_0 + N_{\frac{1}{2}})$.

p. 418, line 9 should read:

. . . Hence if $B_\mu \neq C_\nu$, the . . .

p. 421, line —18 should read:

. . . in the expansion of a, b with

p. 427, line —16 should read:

. . . The elements $e_2, \dots, e_\gamma$ are all orthogonal to e . . .

p. 434, line —5 of text (last part of Eq. (44)) should read:

$$XY = Z$$

ON AN ALGEBRAIC GENERALIZATION OF THE QUANTUM MECHANICAL FORMALISM

Introduction

One of us has shown that the statistical properties of the measurements of a quantum mechanical system assume their simplest form when expressed in terms of a certain hypercomplex algebra which is commutative but not associative.¹ This algebra differs from the non-commutative but associative matrix algebra usually considered in that one is concerned with the commutative expression $\frac{1}{2}(A \times B + B \times A)$ instead of the associative product $A \times B$ of two matrices. It was conjectured that the laws of this commutative algebra would form a suitable starting point for a generalization of the present quantum mechanical theory. The need of such a generalization arises from the (probably) fundamental difficulties resulting when one attempts to apply quantum mechanics to questions in relativistic and nuclear phenomena.

The mathematical foundations underlying this generalization have already been investigated completely² and the results obtained are summarized in Part I of this paper except for the following alterations: 1) We restrict ourselves entirely to commutative systems while in the investigation mentioned the commutative law was partly excluded. Further, it is of basic importance that we assume the "field of coefficients" of the algebra considered to be always real. (In this way we meet an old objection to the quantum mechanical formalism, namely, that it does not confine itself to the real domain. Complex numbers and the like are used here only as technical expedients.) 2) We give the unpublished proof of the deducibility of the law $(a^2b)a = a(ba^2)$ from the law $a^{\mu}a^{\nu} = a^{\mu+\nu}$. 3) We give determinative conditions for the further investigation in Part III of the structure of the " r -number systems"; these conditions appear in (23), (24), and (25) below.

Part II contains a summary discussion of the concepts "trace" and "interchangeability"; the investigation of these concepts has already been begun.²

The essential achievement of this paper is contained in Part III. There the structure of all (formally real, finite, and commutative) r -number systems is fully determined. It is shown that, with but one exception, all such systems are merely matrix algebras. This is noteworthy in itself from the mathematical point of view. But of more importance is the fact that the above mentioned

¹ P. Jordan, *Zschr. f. Phys.*, vol. 80, 1933, p. 285.

² P. Jordan, *Göttinger Nachr.*, 1932, p. 569; 1933, p. 209.

physical considerations demand an algebra essentially more general than matrix algebras. Our results therefore point to the necessity of further generalizations. Although we found *one* r -number system which is not a matrix algebra ($\mathfrak{M}_3^8$, cf. paragraph 22), this system is too narrow for the generalization needed. It may nevertheless be usable as a first step in this direction.

Apart from the possibility of dropping the requirement of commutativity, one first thinks of omitting the condition of finiteness, all the more so since the algebra appearing in quantum mechanics is, as well known, infinite (that is, it has no finite linear basis). It may well happen that new types of algebra will arise with the removal of some of the present restrictions. This possibility is further suggested by the fact that in the ordinary quantum mechanics many important features first appeared in terms of an infinite algebra.³ But it is too early to make conjectures concerning these difficult questions.

However, our procedure can be regarded as a phenomenological basis for the matrix calculus of the quantum theory.

We wish to express here our appreciation to Dr. C. C. Torrance for his valuable assistance in connection with the preparation of this manuscript.

Foundations of the theory

1. We consider an algebra A with elements $a, b, c, \dots$ such that from the elements a and b we can form the sum $a + b$, the difference $a - b$, and the product ab ; also from the element a we can form the element λa , where λ is a real number. (In what follows, greek letters will always denote real numbers.) In our algebra all the usual rules of calculation are to hold (including the commutative law of multiplication) with the single exception of the associative law of multiplication.⁴ Our algebra will have a finite linear base, that is, there will be a finite number of linearly independent elements $a_1, a_2, \dots, a_N$ such that each element of the algebra is linearly dependent upon these a 's. Further, our algebra will be "formally real" in the sense that

$$(I) \quad \text{If } a^2 + b^2 + c^2 + \dots = 0, \text{ then } a = b = c = \dots = 0.$$

We call this algebra an r -number algebra if it has the properties

$$(II) \quad a^\mu a^\nu = a^{\mu + \nu}, \quad (\mu, \nu = 1, 2, \dots)$$

$$(III) \quad (a^2 b) a = a^2 (b a),$$

³ For example, the possibility of satisfying the commutation relations $pq - qp = \frac{h}{2\pi i} 1$, where h is a real number different from zero.

⁴ That is, it is not necessary that $(ab)c = a(bc)$, though we desire that $(\lambda\mu)a = \lambda(\mu a)$ and that $(\lambda a)b = \lambda(ab)$.

⁵ An algebra fulfilling this condition which, in a more general form, is better adapted to non-commutative algebras, has been called by Jordan² *semi-simple*. The present name, i.e., real, (used by E. Artin and I. Schreier, Hamb. Abh., vol. 5, 1926, p. 83) seems to be more suitable.

where a^μ is defined by

$$a^1 = a, \quad a^{\mu+1} = a a^\mu \quad (\mu = 1, 2, \dots).$$

FUNDAMENTAL THEOREM 1. *According to the assumptions made above, including (I), III is a consequence of II and II is a consequence of III.*

The importance of this theorem lies in that (II) must be regarded as a necessary axiom because of the physical meaning of the theory,⁶ whereupon the theorem shows that (III) is also a necessary law of any meaningful generalization of quantum mechanics.

The derivation of (III) from (II) is given at the end of paragraph 4. The derivation of (II) from (III) comes about in the following way:

We use the notation

$$(1) \quad [a, b, c] = (ab)c - a(bc).$$

From the commutative law it follows immediately that

$$(2) \quad [a, b, c] + [c, b, a] = 0,$$

$$(3) \quad [a, b, c] + [b, c, a] + [c, a, b] = 0.$$

If we replace b and a in (III) respectively by u and $\lambda a + \mu b + \nu c$, then an identity results in λ, μ, ν such that the terms proportional to $\lambda\mu\nu$ give

$$(4) \quad [ab, u, c] + [bc, u, a] + [ca, u, b] = 0.$$

Now suppose that it has already been shown that (II) follows from (III) for all μ and ν such that $\mu + \nu \leq N$, where a^μ is determined by the definition given in (III). By (2) and (4) it follows that for $1 \leq \nu \leq N$

$$[a^{\nu+1}, b, a^{N-\nu-1}] = [a^\nu, b, a^{N-\nu}] + [a, b, a^{N-1}].$$

Hence

$$[a^\rho, b, a^{N-\rho}] = \rho[a, b, a^{N-1}] \text{ for } \rho = 1, 2, \dots, N-1,$$

and therefore we have in particular that

$$-[a, b, a^{N-1}] = [a^{N-1}, b, a] = (N-1)[a, b, a^{N-1}],$$

$$[a, b, a^{N-1}] = 0, \quad [a^\rho, b, a^{N-\rho}] = 0.$$

This gives the relation $a^{\rho+1}a^{N-\rho} = a^\rho a^{N-\rho+1} = a^1 a^N = a^{N+1}$ in the case where $b = a$. Hence (II) also holds for $\mu + \nu = N + 1$. Hence (II) holds in general and we have incidentally shown that $[a^\mu, b, a^\nu] = 0$.

In what follows we use (II) as a hypothesis; the further development will show that (III) is a consequence.

2. Before we undertake the actual investigation of A we point out that the mathematical foundation (equivalent to the physical basis, cf. footnote 1) may

⁶ Cf. footnote 1.

also be based on $a \pm b$, λa , and a^μ instead of $a \pm b$, λa , and ab . For then we could define ab as $\frac{1}{2}((a+b)^2 - a^2 - b^2)$ from which it follows automatically that ab is commutative. Accordingly the postulation assumes the following form in this case:

In the algebra A with elements $a, b, c, \dots$, one can form from the elements a and b the elements $a \pm b$, λa (λ a real number), and a^μ ($\mu = 1, 2, \dots, a^1 = a$); all the usual rules of calculation hold for the operations $a \pm b$ and λa ; and A has a finite linear basis. a^μ is a continuous function of a (that is, if we represent all the elements as linear aggregates of the basis elements, then the coefficients of a^μ depend continuously on those of a). A is "formally real", that is,

$$(I') \quad \text{If } a^2 + b^2 + c^2 + \dots = 0, \text{ then } a = b = c = \dots = 0.$$

((I') is the same as (I).) It is always the case that

$$(II') \quad (a^\mu + a^\nu)^2 = a^{2\mu} + a^{2\nu} + 2a^{\mu+\nu}.$$

((II') is equivalent to (II), but it has the definition $a^1 = a$, $a^{\mu+1} = a a^\mu$ as a consequence.) Finally,

$$(*) \quad (a + b + c)^2 + a^2 + b^2 + c^2 = (a + b)^2 + (a + c)^2 + (b + c)^2.$$

((*) is equivalent to $(a + b)c = ac + bc$, but because of the continuity of a^2 (and therefore also of ab), it has the condition $(\lambda a)b = \lambda(ab)$ as a consequence. This completes the list of the desired rules of calculation.—(*) could be replaced by the simpler rule $(a + b)^2 + (a - b)^2 = 2a^2 + 2b^2$.

All of these postulates, except for (*), necessarily arise from the assumed physical conditions.

Part I. Structure of the r-number algebras.

3. We now begin the investigation of the properties of A .

THEOREM 1. *If $a \neq 0$, then $a^\mu \neq 0$.*

Proof: If $a^\nu \neq 0$, then, by (I), $(a^\nu)^2 \neq 0$. Therefore $a^{\nu+1} \neq 0$ since we have $a^{\nu+1}a^{\nu-1} = a^{2\nu} = (a^\nu)^2$. (If $\nu = 1$, the factor $a^{\nu-1}$ is to be omitted.) Hence if $a \neq 0$, then $a^\mu \neq 0$.

We call an element $e = e^2 \neq 0$ a unit element (or idempotent element). If two distinct unit elements e and e' are such that $ee' = 0$, we say that they are mutually orthogonal. If $ee' = 0$, then $e + e'$ is also a unit element (because $(e + e')^2 = e^2 + e'^2 + 2ee' = e + e'$ and $e(e + e') = e^2 + ee' = e \neq 0$ so that $e + e' \neq 0$).

THEOREM 2. *Corresponding to each $a \neq 0$ there exists a unit element $e = e(a)$ (where e is a polynomial in a , and $e = e^2 \neq 0$) such that, for $\mu \geq 2$, $ea^\mu = a^\mu$.*

Proof: Since the basis of A is finite, only a finite number of the elements $a, a^2, \dots$ are linearly independent. Suppose that a^μ depends linearly on $a, \dots, a^{\mu-1}$ so that

$$a^\mu + \lambda_{\mu-1} a^{\mu-1} + \dots + \lambda_0 a^0 = 0.$$

Since $a \neq 0$ and $a^\mu \neq 0$, it follows that $\lambda_\rho \neq 0$ and $\rho < \mu$ if ρ is given its largest possible value. The condition $\rho > 2$ would imply that

$$a^\mu + \lambda_{\mu-1}a^{\mu-1} + \dots + \lambda_\rho a^\rho = (a^{\mu-2} + \lambda_{\mu-1}a^{\mu-3} + \dots + \lambda_\rho a^{\rho-2})a^2 = 0,$$

$$(a^{\mu-2} + \lambda_{\mu-1}a^{\mu-3} + \dots + \lambda_\rho a^{\rho-2})^2 a^2 = 0,$$

so that, by Theorem 1,

$$(a^{\mu-2} + \lambda_{\mu-1}a^{\mu-3} + \dots + \lambda_\rho a^{\rho-2})a = a^{\mu-1} + \lambda_{\mu-1}a^{\mu-2} + \dots + \lambda_\rho a^{\rho-1} = 0.$$

Therefore if μ is given its smallest possible value it follows that $\rho = 1$ or 2 . Now put

$$(5) \quad e = e(a) = \frac{(a^{\mu-\rho} + \lambda_{\mu-1}a^{\mu-\rho-1} + \dots + \lambda_{\rho+1}a)^\rho}{(-\lambda_\rho)^\rho}.$$

Then

$$\frac{a^{\mu-\rho} + \lambda_{\mu-1}a^{\mu-\rho-1} + \dots + \lambda_{\rho+1}a}{-\lambda_\rho} a^\rho = a^\rho.$$

Hence $ea^\rho = a^\rho$, and $ea^{\mu'} = a^{\mu'}$ if $\mu' \geq \rho$, and therefore *a fortiori* if $\mu' \geq 2$. Hence $e^2 = e e(a) = e$, and since $ea^2 = a^2 \neq 0$, it follows that $e \neq 0$, i.e., e is a unit element. (Note that, having dealt in this proof only with powers and polynomials involving one element a , we could use the associative law because of (II).)

The simplest special case of (II) is $a^2a^2 = a^4$, that is, $a^2a^2 = (a^2a)a$. If we now replace a by $\lambda a + \mu b$ and compare the terms proportional to $\lambda^3\mu$ and $\lambda^2\mu^2$ respectively, then

$$(6) \quad 4(ab)a^2 = a^3b + a(a^2b) + 2a(a(ab)),$$

$$(7) \quad 4(ab)^2 + 2a^2b^2 = a(ab^2) + b(ba^2) + 2a(b(ab)) + 2b(a(ab)).$$

If in (6) we replace a by a unit element e and replace b by x , then

$$(8) \quad 2e(e(x)) - 3e(ex) + ex = 0.$$

THEOREM 3. *If e is a unit element and if $N_\lambda = N_\lambda(e)$ is the set of x 's such that $ex = \lambda x$, then N_λ is obviously a linear manifold, but if $\lambda \neq 0, \frac{1}{2}$, and 1 , then N_λ consists of zero alone.*

Each x can be represented in the form

$$(9) \quad x = x_0 + x_{\frac{1}{2}} + x_1, \quad x_0 \in N_0, \quad x_{\frac{1}{2}} \in N_{\frac{1}{2}}, \quad x_1 \in N_1,$$

in one and only one way.

Proof: If $x \in N_\lambda$, then, by (8), $(2\lambda^3 - 3\lambda^2 + \lambda)x = 0$, and if $\lambda \neq 0, \frac{1}{2}$, and 1 , then $x = 0$. If $x \in N_\lambda$, then $2e(ex) - 3ex + x = (2\lambda^2 - 3\lambda + 1)x$. For $\lambda = 0$ this expression reduces to x and for $\lambda = \frac{1}{2}$ or 1 it reduces to 0 . Hence, by (9),

$$(10) \quad \begin{cases} x_0 = 2e(ex) - 3ex + x, \\ x_{\frac{1}{2}} = -4e(ex) + 4ex, \\ x_1 = 2e(ex) - ex, \end{cases}$$

where the latter two formulas are obtained in a similar way. Thus equation (9) can have no solution other than (10). But (10) is a solution, for the equations $ex_0 = 0$, $ex_{\frac{1}{2}} - \frac{1}{2}x_{\frac{1}{2}} = 0$, and $ex_1 - x_1 = 0$ are equivalent to (8), and it is evident that $x = x_0 + x_{\frac{1}{2}} + x_1$.

THEOREM 4. *If $a \in N_0$ and $b \in N_0$, then $ab \in N_0$; if $a \in N_1$ and $b \in N_1$, then $ab \in N_1$; if $a \in N_0$ and $b \in N_1$, then $ab = 0$. Thus N_0 and N_1 are mutually orthogonal sub-algebras of A . If $a \in N_{\frac{1}{2}}$ and $b \in N_{\frac{1}{2}}$, then $ab \in (N_0 + N_1)$; if $a \in N_{\frac{1}{2}}$ and $b \in N_0$, then⁷ $ab \in (N_0 + N_{\frac{1}{2}})$; if $a \in N_{\frac{1}{2}}$ and $b \in N_1$, then⁷ $ab \in (N_{\frac{1}{2}} + N_1)$.*

Proof: in (7) we set $a = e$ and replace b by $\lambda a + \mu b$. If we consider the terms proportional to $\lambda\mu$, then

$$8(ea)(eb) + 4e(ab) - 2e(e(ab)) - (ea)b - (eb)a \\ - 2e(a(eb)) - 2e(b(ea)) - 2a(e(eb)) - 2b(e(ea)) = 0,$$

and if $a \in N_\rho$ and $b \in N_\sigma$, then

$$- 2e(e(ab)) + (4 - 2\rho - 2\sigma)e(ab) + (8\rho\sigma - 2\rho^2 - 2\sigma^2 - \rho - \sigma)ab = 0.$$

If we resolve ab according to Theorem 3, then

$$(8\rho\sigma - 2\rho^2 - 2\sigma^2 - \rho - \sigma)(ab)_0 + (8\rho\sigma - 2\rho^2 - 2\sigma^2 - 2\rho - 2\sigma + \frac{3}{2})(ab)_{\frac{1}{2}} \\ + (8\rho\sigma - 2\rho^2 - 2\sigma^2 - 3\rho - 3\sigma + 2)(ab)_1 = 0.$$

Since this may be regarded as a resolution of zero in the sense of Theorem 3, each of the three terms separately vanish. Hence if one of the three conditions

$$8\rho\sigma - 2\rho^2 - 2\sigma^2 - \rho - \sigma \neq 0, \quad 8\rho\sigma - 2\rho^2 - 2\sigma^2 - 2\rho - 2\sigma + \frac{3}{2} \neq 0, \\ 8\rho\sigma - 2\rho^2 - 2\sigma^2 - 3\rho - 3\sigma + 2 \neq 0$$

is satisfied, the corresponding equation $(ab)_0 = 0$, $(ab)_{\frac{1}{2}} = 0$, $(ab)_1 = 0$ must be true. Substitution of $\rho, \sigma = 0, \frac{1}{2}$, and 1 leads to the desired results.

COROLLARY. *If $a \in N_{\frac{1}{2}}$, then, by Theorem 4, $a^2 \in (N_0 + N_1)$; yet were it the case that $a^2 \in N_0$ or $a^2 \in N_1$, then a would be zero.*

Proof: It would follow that $ea = \frac{1}{2}a$ and $ea^2 = \lambda a^2$ with $\lambda = 0$ or 1. Hence if $b = e$, then (6) becomes

$$2a^3 = ea^3 + \lambda a^3 + a^3,$$

so that $ea^3 = (1 - \lambda)a^3$ and $a^3 \in N_{1-\lambda}$. By Theorem 4, it follows that $a^5 = a^3a^2 = 0$ ($a^2 \in N_\lambda$, $a^3 \in N_{1-\lambda}$), so that $a = 0$.

THEOREM 5. *There exists one modulus 1, that is, one unit element $e = 1$ such that $1a = a$ for every a .*

Proof: We select a unit element e so that the resulting dimension δ of the algebra $N_0(e)$ is as small as possible. Suppose $N_0(e)$ contained an element $a \neq 0$. Then, by Theorem 2, it would contain a unit element e' . Since $e' \in N_0(e)$, it

⁷ We shall later (Theorem 7) restrict this to $ab \in N_{\frac{1}{2}}$.

follows that $ee' = 0$. Hence $e + e'$ is also a unit element. Since $(e + e')e = e$, e belongs to $N_1(e + e')$, and if $a \in N_0(e + e')$ then, by Theorem 4, $ea = 0$ and $a \in N_0(e)$. Therefore $N_0(e + e') \subseteq N_0(e)$. $e' \in N_0(e)$, but e' does not belong to $N_0(e + e')$ since $(e + e')e' = e' \neq 0$. Therefore $N_0(e + e') \neq N_0(e)$. Hence the dimension of $N_0(e + e')$ is less than δ , contrary to assumption. Consequently $N_0(e)$ consists of 0 alone.

If $a \in N_1(e)$, then, by Theorem 4, $a^2 \in [N_0(e) + N_1(e)]$. By the above argument $a^2 \in N_1(e)$, and by the Corollary, $a = 0$. Hence $N_1(e)$ consists of 0 alone.

In the notation of (9) it follows that x always equals x_1 , hence that $x \in N_1(e)$, and hence that $ex = x$. Therefore e is the modulus.

4. Because of the existence of the modulus we may restate in more complete form the considerations which led up to Theorem 2. Among the elements 1, $a, a^2, \dots$ linear dependence must occur so that there exists a relation of the form

$$(11) \quad f(a) = a^\mu + \lambda_{\mu-1}a^{\mu-1} + \dots + \lambda_1a + \lambda_01 = 0.$$

If this relation is such that μ has the smallest possible value, then we call (11) the *characteristic equation* and we call μ the *order* of a . These notions are uniquely determined, for if there existed two equations $f(a) = 0$ and $g(a) = 0$ of this form, then $f(a) - g(a) = 0$ would be an equation of lower degree unless the polynomials f and g were identical.

Let ξ be a variable whose values are real numbers. Suppose the equation $f(\xi) = 0$ had a complex root $\xi = \rho + i\sigma$ with $\sigma \neq 0$. Then it would also have the root $\xi = \rho - i\sigma$. Hence $f(\xi)$ would be divisible by the polynomial

$$(\xi - (\rho + i\sigma))(\xi - (\rho - i\sigma)) = (\xi - \rho)^2 + \sigma^2.$$

Again, if $f(\xi) = 0$ had a multiple real root $\xi = \rho$, then $f(\xi)$ would again be divisible by $(\xi - \rho)^2 + \sigma^2$, where $\sigma = 0$. In either case $f(\xi)$ is of the form

$$f(\xi) = ((\xi - \rho)^2 + \sigma^2)g(\xi).$$

If $f(a) = 0$, then $f(a)g(a) = 0$, and hence $((a - \rho)g(a))^2 + (\sigma g(a))^2 = 0$. It follows by (I) that $(a - \rho)g(a) = 0$. But this is impossible since $(\xi - \rho)g(\xi)$ is a polynomial of the same form as $f(\xi)$ but of lower degree.

It follows that $f(\xi) = 0$ has μ distinct simple real roots $\xi_1, \xi_2, \dots, \xi_\mu$ which we call the *proper values* of a .

If two polynomials $g(\xi)$ and $h(\xi)$ have the same value at each of the points $\xi = \xi_1, \xi_2, \dots, \xi_\mu$, then $g(\xi) - h(\xi)$ is divisible by $f(\xi) = (\xi - \xi_1) \dots (\xi - \xi_\mu)$ so that $g(\xi) - h(\xi) = f(\xi)k(\xi)$. Hence

$$g(a) - h(a) = f(a)k(a) = 0 \quad \text{and} \quad g(a) = h(a).$$

THEOREM 6. *Each element a may be represented as a linear combination of a set of pairwise orthogonal unit elements whose sum is the modulus; the coefficients in this representation are the proper values of a .*

Proof: Let $p_\nu(\xi) = \frac{f(\xi)}{(\xi - \xi_\nu)f'(\xi_\nu)}$ ($\nu = 1, 2, \dots, \mu$) be that polynomial whose value is zero when ξ has any of the values $\xi_1, \xi_2, \dots, \xi_\mu$ except ξ_ν , in which case $p_\nu(\xi)$ has the value unity. Since $p_\nu(\xi)$ is not identically zero and since the degree of $p_\nu(\xi)$ is less than μ , it follows that $p_\nu(a) \neq 0$. Let

$$(12) \quad e_\nu = p_\nu(a) \quad (\nu = 1, 2, \dots, \mu).$$

For each $\xi = \xi_1, \xi_2, \dots, \xi_\mu$, $p_\nu(\xi)^2 = p_\nu(\xi)$ and $p_\nu(\xi)p_{\nu'}(\xi) = 0$, where $\nu \neq \nu'$; also $p_1(\xi) + \dots + p_\mu(\xi) = 1$ and $\xi_1 p_1(\xi) + \dots + \xi_\mu p_\mu(\xi) = \xi$. It follows from the last remark before Theorem 6 that

$$(13) \quad \begin{cases} e_\nu^2 = e_\nu \neq 0, e_\nu e_{\nu'} = 0 \text{ for } \nu \neq \nu', \\ e_1 + \dots + e_\mu = 1, \\ \xi_1 e_1 + \dots + \xi_\mu e_\mu = a. \end{cases}$$

This completes the proof of all parts of the theorem.

Since the relation $g(\xi_1)p_1(\xi) + \dots + g(\xi_\mu)p_\mu(\xi) = g(\xi)$ holds for any polynomial $g(\xi)$ for each $\xi = \xi_1, \xi_2, \dots, \xi_\mu$, then

$$(14) \quad g(\xi_1)e_1 + \dots + g(\xi_\mu)e_\mu = g(a).$$

If $\xi_\nu \neq 0$, we can find a polynomial $h(\xi)$ which is zero for $\xi = 0, \xi_1, \xi_2, \dots$, and ξ_μ except ξ_ν , and which has the value unity for $\xi = \xi_\nu$. Then by (14), $h(a) = e_\nu$ and $h(\xi)$ contains no constant term. Hence $h(a)$ contains no term 1 but only terms $a, a^2, \dots$. This proves the following

COROLLARY. *If a belongs to a sub-algebra B of A , then each of its e_ν whose $\xi_\nu \neq 0$ belongs to B .*

THEOREM 7. *Theorem 4 may be modified as follows: if $a \in N_{\frac{1}{2}}$ and if either $b \in N_0$ or $b \in N_1$, then $ab \in N_{\frac{1}{2}}$.*

Proof: By Theorem 6 and its corollary it is sufficient, since N_0 and N_1 are algebras, to consider only unit elements $b = e'$. $1 - e$ is a unit element, for $(1 - e)^2 = 1 - e$, and if $1 - e$ were zero, then it would be the case that $e = 1$ and $N_{\frac{1}{2}} = 0$, so that no proof would be necessary. Since the substitution of $(1 - e)$ for e carries $N_0, N_{\frac{1}{2}}$, and N_1 over into $N_1, N_{\frac{1}{2}}$, and N_0 respectively, it is sufficient to consider only those elements $b \in N_0$; for these elements, $ee' = 0$. (In this argument N_λ has really signified $N_\lambda(e)$.)

Since e and e' are orthogonal it follows that $e + e'$ is a unit element and according to (9) we are able to resolve a with respect to $e + e'$ as follows: $a = a_0 + a_{\frac{1}{2}} + a_1$. Since $a \in N_{\frac{1}{2}}(e)$ and $ea = \frac{1}{2}a$, we have

$$ea_0 + ea_{\frac{1}{2}} + ea_1 = \frac{1}{2}a_0 + \frac{1}{2}a_{\frac{1}{2}} + \frac{1}{2}a_1.$$

By application of Theorem 4 to $e + e'$ it follows that

$$ea_0 = 0, \quad ea_{\frac{1}{2}} \in [N_0(e + e') + N_{\frac{1}{2}}(e + e')], \quad ea_1 \in N_1(e + e')$$

since $(e + e')e = e$ and $e \in N_1(e + e')$. From the right hand side of the above equation we have that

$$(\frac{1}{2}a_0 + \frac{1}{2}a_1) \in [N_0(e + e') + N_1(e + e')], \quad \frac{1}{2}a_1 \in N_1(e + e').$$

Because of the uniqueness of the resolution (9) it results that

$$ea_{\frac{1}{2}} = \frac{1}{2}a_0 + \frac{1}{2}a_1, \quad ea_1 = \frac{1}{2}a_1.$$

By definition,

$$(e + e')a_0 = 0, \quad (e + e')a_{\frac{1}{2}} = \frac{1}{2}a_{\frac{1}{2}}, \quad (e + e')a_1 = a_1,$$

so that

$$e'a_0 = 0, \quad e'a_{\frac{1}{2}} = -\frac{1}{2}a_0, \quad e'a_1 = \frac{1}{2}a_1.$$

Hence $e'(e'a_{\frac{1}{2}}) = 0$. Resolving $a_{\frac{1}{2}}$ by (9) with respect to e' shows that this implies $a_{\frac{1}{2}} \in N_0(e')$. Therefore $-\frac{1}{2}a_0 = e'a_{\frac{1}{2}} = 0$ and $a_0 = 0$. Consequently

$$e'a = e'a_0 + e'a_{\frac{1}{2}} + e'a_1 = 0 + 0 + \frac{1}{2}a_1 = \frac{1}{2}a_1.$$

Since $ea_1 = \frac{1}{2}a_1$, $a_1 \in N_{\frac{1}{2}}(e)$, and $e'a \in N_{\frac{1}{2}}(e)$, we have finally that $ab \in N_{\frac{1}{2}}(e)$.

Theorem 6 and Formula (14) give us the means to prove (III), that is, that $[a^2, b, a] = 0$. If we resolve a by Theorem 6 then, by (14), it is sufficient to prove the relation $[e_{\nu}, b, e_{\nu'}] = 0$ for all ν and ν' . If $\nu = \nu'$ this relation follows from (2), but if $\nu \neq \nu'$, then $e_{\nu}e_{\nu'} = 0$ and it remains to show that

$$(15) \quad e(e'b) = e'(eb)$$

if e and e' are orthogonal unit elements. By Theorem 3 it suffices to consider the cases $b \in N_0(e)$, $b \in N_{\frac{1}{2}}(e)$, and $b \in N_1(e)$. $ee' = 0$ implies $e' \in N_0(e)$, and so theorems 4 and 7 prove the statement.

5. In accordance with the usual general terminology, the algebra A is called *irreducible* if it contains no invariant subalgebra B , that is, if no subalgebra B (except A itself and 0) is such that if $a \in A$ and $b \in B$, then $ab \in B$. A is called a *direct sum* if it is the sum of two orthogonal subalgebras B and C , that is, if every $a \in A$ can be written in the form $a = b + c$, where $b \in B$ and $c \in C$ and where B and C are such that if $b \in B$ and $c \in C$, then $bc = 0$ (neither B nor C being A or 0).⁸ Every r -number algebra is completely reducible, that is, the properties of being reducible and of being a direct sum are equivalent. For it is clear that B is an invariant subalgebra if A is the direct sum of B and C . Conversely, if B is an invariant subalgebra then it has a modulus 1^* , and from its invariance it follows immediately that $N_1(1^*) = B$, $N_{\frac{1}{2}}(1^*) \subset B$, so that $N_{\frac{1}{2}}(1^*) = 0$. Consequently A is the direct sum of B and C if $N_0(1^*) = C$.

⁸ If b belongs to B and C simultaneously, then $b^2 = 0$ and $b = 0$. Hence the resolution $a = b + c$ is unique, where $b \in B$ and $c \in C$.

We can represent every r -number algebra A as the direct sum of finitely many irreducible subalgebras $B_1, \dots, B_k$ (all $\neq 0$), that is, each $a \in A$ is

$$(16) \quad a = b_1 + \dots + b_k, \quad b_1 \in B_1, \dots, b_k \in B_k,$$

where if $b_i \in B_i$ and $b_j \in B_j$ for $i \neq j$, then $b_i b_j = 0$.⁹ Moreover, $B_1, \dots, B_k$ are uniquely determined except for their order. For let $C_1, \dots, C_h$ be a second resolution of this kind. Since B_μ and C_ν are invariant subalgebras of A , so also is their intersection which is likewise an invariant subalgebra of B_μ and C_ν . Since both B_μ and C_ν are irreducible, their intersection must be 0 or else it must equal B_μ and C_ν . Hence if $B_\mu = C_\nu$, the intersection of B_μ and C_ν is zero. The sum of all the C_ν is A which in turn contains B_μ . Thus not all of these intersections can be zero. Hence for each μ there is a ν such that $B_\mu = C_\nu$. Because the C_ν are pairwise orthogonal there is only one such ν . Similarly for each ν there is one and only one μ such that $B_\mu = C_\nu$. Hence $k = h$ and the B_μ go over by a permutation into the C_ν .

The $B_1, \dots, B_k$ are the *irreducible components* of A . It is apparent from the above discussion that it is sufficient to investigate the irreducible algebras A since the other algebras may be constructed from these in the simple manner just indicated. Consequently we shall confine our interest merely to the irreducible algebras, although we shall word our theorems for all algebras wherever the proofs do not make use of irreducibility.

6. We now examine systems of pairwise orthogonal unit elements $e_1, \dots, e_\gamma$ ($e_\rho^2 = e_\rho \neq 0, e_\rho e_\sigma = 0$ for $\rho \neq \sigma$).

THEOREM 8. Let $e_1, \dots, e_\gamma$ be a system of pairwise orthogonal unit elements with the sum 1. Let the set of all x such that¹⁰

$$(17) \quad e_\tau x = \frac{1}{2}(\delta_{\rho\tau} + \delta_{\sigma\tau})x \quad (\tau = 1, \dots, \gamma)$$

be denoted by $M^{\rho\sigma}$, where $\rho, \sigma = 1, \dots, \gamma$.¹¹ $M^{\rho\sigma}$ is obviously a linear manifold.

Every x can be represented in the form

$$(18) \quad x = \sum_{\rho \leq \sigma} x_{\rho\sigma}, \quad x_{\rho\sigma} \in M^{\rho\sigma},$$

in one and only one way.

Proof: It follows from (17) and (18) that

$$(19) \quad \begin{cases} 4e_\rho(e_\sigma x) = 4e_\sigma(e_\rho x) = x_{\rho\sigma} & (\rho \neq \sigma), \\ 2e_\rho(e_\rho x) - e_\rho x = x_{\rho\rho}. \end{cases}$$

⁹ Consequently the resolution (16) is unique. Cf. footnote 8.

¹⁰ $\delta_{\mu\nu}$ is the Kronecker symbol, that is, $\delta_{\mu\nu} = \begin{cases} 1 & \text{for } \mu = \nu, \\ 0 & \text{for } \mu \neq \nu. \end{cases}$

¹¹ It is obvious that $M^{\rho\sigma} = M^{\sigma\rho}$, so that for the most part we shall be able to limit ourselves to the consideration of $M^{\rho\sigma}$ for $\rho \leq \sigma$.

Therefore (18) can have no solution other than (19) and, as we shall show, (19) is a solution of (18). ((15) shows that (19) is self-consistent.) Comparison of (19) with (10) shows that $x_{\rho\rho}$ belongs to $N_1(e_\rho)$. For $\tau \neq \rho$, $e_\tau e_\rho = 0$ so that these e_τ belong to $N_0(e_\rho)$. Hence for $\tau \neq \rho$, $e_\rho x_{\rho\rho} = x_{\rho\rho}$ and $e_\tau x_{\rho\rho} = 0$. Hence by (17), $x_{\rho\rho} \in M^{\rho\rho}$. If in (4) we set $a = b = e_\rho$, $c = x$, and $u = e_\sigma$, then for $\rho \neq \sigma$ (since $e_\rho^2 = e_\rho$ and $e_\rho e_\sigma = 0$)

$$[e_\rho, e_\sigma, x] + 2[e_\rho x, e_\sigma, e_\rho] = 0,$$

$$-e_\rho(e_\sigma x) + 2e_\rho(e_\sigma(e_\rho x)) = 0,$$

$$e_\rho(e_\sigma(e_\rho x)) = \frac{1}{2}e_\rho(e_\sigma x), \quad e_\rho x_{\rho\sigma} = \frac{1}{2}x_{\rho\sigma}.$$

By symmetry it follows that $e_\sigma x_{\rho\sigma} = \frac{1}{2}x_{\rho\sigma}$. Hence $(e_\rho + e_\sigma)x_{\rho\sigma} = x_{\rho\sigma}$ and $x_{\rho\sigma} \in N_1(e_\rho + e_\sigma)$. If $\tau \neq \rho$ and σ , then $(e_\rho + e_\sigma)e_\tau = 0$ and $e_\tau \in N_0(e_\rho + e_\sigma)$ so that $e_\tau x_{\rho\sigma} = 0$. By (17), $x_{\rho\sigma} \in M^{\rho\sigma}$. Finally, (since $\sum_\rho e_\rho = 1$)

$$\begin{aligned} \sum_{\rho \leq \sigma} x_{\rho\sigma} &= \sum_\rho (2e_\rho(e_\rho x) - e_\rho x) + \sum_{\rho < \sigma} (2e_\rho(e_\sigma x) + 2e_\sigma(e_\rho x)) \\ &= 2 \sum_{\rho, \sigma} e_\rho(e_\sigma x) - \sum_\rho e_\rho x = 2x - x = x. \end{aligned}$$

This completes the proof of all parts of the theorem.

THEOREM 9. Suppose $a \in M^{\rho\sigma}$ and $b \in M^{\pi\tau}$. Then: if both ρ and σ are distinct from both π and τ , then $ab = 0$; if ρ, σ and π, τ have one common index, say $\rho = \pi$ but $\sigma \neq \tau$, then $ab \in M^{\sigma\tau}$; if ρ, σ are the same pair as π, τ (except possibly for order; cf. footnote 11), say $\rho = \pi$ and $\sigma = \tau$, then $ab \in (M^{\rho\rho} + M^{\sigma\sigma})$ for $\rho \neq \sigma$ and $ab \in M^{\rho\rho}$ for $\rho = \sigma$.

Proof: We set

$$e = \begin{cases} e_\rho + e_\sigma & \text{for } \rho \neq \sigma, \\ e_\rho & \text{for } \rho = \sigma, \end{cases} \quad e' = \begin{cases} e_\pi + e_\tau & \text{for } \pi \neq \tau, \\ e_\pi & \text{for } \pi = \tau. \end{cases}$$

By (17) it follows that $ea = a$ and $e'b = b$ so that $a \in N_1(e)$ and $b \in N_1(e')$. If both ρ and σ are distinct from both π and τ , then it is obvious that $ee' = 0$ and $e' \in N_0(e)$. Hence, by Theorem 4, $e'a = 0$, $a \in N_0(e')$, and $ab = 0$.

Suppose that σ, τ , and $\rho = \pi$ are three distinct numbers. By Theorem 4, $B = N_1(e_\rho + e_\sigma + e_\tau)$ is an algebra and its modulus is $e_\rho + e_\sigma + e_\tau$. If we resolve the condition $(e_\rho + e_\sigma + e_\tau)x = x$ according to (17) and (18), the result shows that x can contain only those addends $x_{\mu\nu}$ in which $\mu, \nu = \rho, \sigma, \tau$. Hence

$$B = N_1(e_\rho + e_\sigma + e_\tau) = \sum_{\mu, \nu = \rho, \sigma, \tau} M^{\mu\nu}. \quad \text{In a similar way it is apparent}$$

$$\text{that the relations } N_1(e_\rho + e_\sigma) = \sum_{\mu, \nu = \rho, \sigma} M^{\mu\nu}, \quad N_1(e_\rho + e_\tau) = \sum_{\mu = \rho, \sigma} M^{\mu\tau}, \text{ and}$$

$N_0(e_\rho + e_\sigma) = M^{\tau\tau}$ hold in B . So $a \in N_1$ and $b \in N_1$ and, by Theorem 7, $ab \in N_1$, where $N_1 = M^{\rho\rho} + M^{\sigma\sigma}$. By symmetry it follows that $ab \in (M^{\rho\sigma} + M^{\sigma\rho})$. Because of the uniqueness of the resolution (18), $ab \in M^{\sigma\tau}$.

If $\rho = \pi = \sigma \neq \tau$, we reach the same result by considering $B = N_1(e_\rho + e_\tau)$ and $N_1(e_\rho)$, $N_1(e_\tau)$, and $N_0(e_\rho)$. By symmetry the same result holds for $\rho = \pi = \tau \neq \sigma$. This accounts for all the cases where $\rho = \pi$ and $\sigma \neq \tau$.

If $\rho = \pi$ and $\sigma = \tau$, then $e = e'$. It follows that $a \in N_1(e)$ and $b \in N_1(e)$ so that $ab \in N_1(e)$. If $\rho = \sigma$, then $ab \in N_1(e_\rho)$, where $N_1(e_\rho) = M^{\rho\rho}$. If $\rho \neq \sigma$, then $ab \in N_1(e_\rho + e_\sigma)$, where $N_1(e_\rho + e_\sigma) = \sum_{\mu, \nu = \rho, \sigma} M^{\mu\nu} = M^{\rho\rho} + M^{\sigma\sigma} + M^{\rho\sigma}$.

The relations $N_1(e_\rho) = M^{\rho\rho}$, $N_1(e_\sigma) = M^{\sigma\sigma}$, and $N_0(e_\rho) = M^{\sigma\sigma}$ (note (17) and (18)) hold in $B = N_1(e_\rho + e_\sigma)$. Hence $a \in N_1$, $b \in N_1$, and, by Theorem 4, $ab \in (N_0 + N_1)$, where $N_0 + N_1 = M^{\rho\rho} + M^{\sigma\sigma}$. This completes the proof of all parts of the theorem.

A unit element e is called *unresolvable* if it can not be represented as the sum of two orthogonal unit elements.

THEOREM 10. *Every unit element can be represented as the sum of pairwise orthogonal unresolvable unit elements.*

Proof: There exist representations of e as a sum of pairwise orthogonal unit elements:

$$e_1 + \cdots + e_\gamma = e.$$

For example, $e_1 = e$, $\gamma = 1$. Since $e_1, \cdots, e_\gamma$ are linearly independent (if $\lambda_1 e_1 + \cdots + \lambda_\gamma e_\gamma = 0$, then $e_\mu(\lambda_1 e_1 + \cdots + \lambda_\gamma e_\gamma) = \lambda_\mu e_\mu = 0$ and $\lambda_\mu = 0$ for all $\mu = 1, \cdots, \gamma$), then $\gamma \leq N$, where N is the dimension of A . Let γ be as large as possible. Suppose an e_μ ($\mu = 1, \cdots, \gamma$) were resolvable, $e_\mu = e'_\mu + e''_\mu$. Then we would have

$$e_1 + \cdots + e'_\mu + e''_\mu + \cdots + e_\gamma = e.$$

Since $e_\mu e'_\mu = (e'_\mu + e''_\mu)e'_\mu = e'_\mu$, therefore $e'_\mu \in N_1(e_\mu)$. Likewise $e''_\mu \in N_1(e_\mu)$. For $\mu \neq \nu$, $e_\mu e_\nu = 0$ and $e_\nu \in N_0(e_\mu)$ so that $e'_\mu e_\nu = e''_\mu e_\nu = 0$. Thus the $\gamma + 1$ elements $e_\mu, \cdots, e'_\mu, e''_\mu, \cdots, e_\gamma$ are pairwise orthogonal. This contradicts our original assumption and every e_μ must be unresolvable.

THEOREM 11. *A unit element e is unresolvable when and only when the dimension of $N_1(e)$ is equal to unity.*

Proof: If e is resolvable, $e = e' + e''$, then $ee' = (e' + e'')e' = e'$ and $e' \in N_1(e)$. Likewise $e'' \in N_1(e)$, and e' and e'' are linearly independent. Hence the dimension is greater than unity.

Conversely, suppose that e is unresolvable. If $a \in N_1(e)$ then, in the resolution of a according to (13), each e_ν (for which $\xi_\nu \neq 0$) belongs to $N_1(e)$ by the corollary of Theorem 6. Hence $ee_\nu = e_\nu$. Since $e_\nu + (e - e_\nu) = e$, $(e - e_\nu)^2 = e - e_\nu$, and $e_\nu(e - e_\nu) = 0$, it follows that if $(e - e_\nu)$ were different from zero then e would be resolvable. Hence $e_\nu = e$. Hence there is at most one $\xi_\nu \neq 0$ so that $a = 0$ or $a = \xi_\nu e$. Therefore $N_1(e)$ contains only the elements λe and the dimension is unity.

On the basis of Theorem 10 we now assume a resolution of 1 in pairwise orthogonal unresolvable unit elements:

$$e_1 + \cdots + e_\gamma = 1,$$

and we form the corresponding $M^{\rho\sigma}$ in accordance with Theorem 8. Since $N_1(e_\rho) = M^{\rho\rho}$, then, by Theorem 11, each $M^{\rho\rho}$ has the dimension unity.

THEOREM 12. *Under the above assumptions $a^2 = \alpha(e_\rho + e_\sigma)$, where a runs through $M^{\rho\sigma}$ and $\rho \neq \sigma$.*

Proof: By Theorem 9 it follows that $a^2 \in (M^{\rho\rho} + M^{\sigma\sigma})$ and, by Theorem 11, $M^{\rho\rho}$ consists merely of the elements αe_ρ . Likewise $M^{\sigma\sigma}$ consists of the elements βe_σ , so that $a^2 = \alpha e_\rho + \beta e_\sigma$. By (III) and (17),

$$(a^2 e_\rho) a = a^2 (e_\rho a), \quad \alpha e_\rho \cdot a = (\alpha e_\rho + \beta e_\sigma) \cdot (e_\rho a),$$

$$\alpha \cdot \frac{1}{2} a = (\alpha + \beta) \cdot \frac{1}{2} \cdot \frac{1}{2} a, \quad \frac{1}{4} (\alpha - \beta) \cdot a = 0,$$

so that if $a \neq 0$ then $\frac{1}{4} (\alpha - \beta) = 0$ and $\alpha = \beta$. If $a = 0$ then it follows immediately that $\alpha = \beta = 0$.

THEOREM 13. *Under the above assumptions there exists a linear basis $s_1, \dots, s_x$ (or if it is desired to show the dependency of the basis on ρ and σ , $s_1^{\rho\sigma}, \dots, s_{x_{\rho\sigma}}^{\rho\sigma}$) of $M^{\rho\sigma}$ such that the following multiplication rules are valid (see footnote 10):*

$$(20) \quad s_\mu s_\nu = \delta_{\mu\nu} (e_\mu + e_\nu) \quad (\mu, \nu = 1, \dots, x).$$

This basis is uniquely determined to within a (real) orthogonal transformation

$$(21) \quad s'_\mu = \sum_{\nu=1}^x \omega_{\mu\nu} s_\nu \quad (\mu = 1, \dots, x),$$

where the matrix $\|\omega_{\mu\nu}\|$ is orthogonal.

Proof: Since $ab = \frac{1}{2}((a+b)^2 - a^2 - b^2)$ it follows from Theorem 12 that ab likewise has the form $\alpha(e_\rho + e_\sigma)$. Therefore $\alpha = \alpha(a, b)$ is obviously linear in a and in b , is symmetric in a and b , and is consequently a symmetric bilinear form in a and b (that is, in the coefficients appearing in the expansion of ab with respect to $s_1, \dots, s_x$). If an $\alpha(a, a)$ were not greater than zero, then it would be the case that

$$a^2 + (\sqrt{-\alpha(a, a)} (e_\rho + e_\sigma))^2 = \alpha(a, a) (e_\rho + e_\sigma) - \alpha(a, a) (e_\rho + e_\sigma) = 0$$

so that, by (I), $a = 0$. Hence if $a \neq 0$ then $\alpha(a, a) > 0$. Therefore the quadratic form $\alpha(a, a)$ is positive definite. Consequently $\alpha(a, b)$ can be transformed into the unit form by a linear transformation. Hence it is possible to select the basis of $M^{\rho\sigma}$ so that $\alpha(a, b)$ is the unit form of the expansion coefficients of a and b . $\alpha(s_\mu, s_\nu) = \delta_{\mu\nu}$ and (20) is proved.

It is clear that (20) is invariant under (21).

Formulas (17) and (20) together with the first and second multiplication rules of Theorem 9 (where we interchange ρ and σ) lead to the following multiplication table, providing that e_ρ and $s_\mu^{\rho\sigma}$ are selected in accordance with Theorems 12 and 13:

THEOREM 14. *Under the above assumptions it follows that e_ρ ($\rho = 1, \dots, \gamma$) and $s_\mu^{\rho\sigma}$ ($\rho, \sigma = 1, \dots, \gamma; \mu = 1, \dots, x_{\rho\sigma}$, where $\rho \neq \sigma$ and the pairs ρ, σ and σ, ρ are to be identified; see footnote 11) is a linear basis of A for which the following multiplication rules hold:*

$$(22) \quad \left\{ \begin{array}{l} \text{(a)} \quad e_\rho e_\sigma = \begin{cases} e_\rho \text{ for } \rho = \sigma, \\ 0 \text{ for } \rho \neq \sigma, \end{cases} \\ \text{(b)} \quad e_\rho s_\mu^{\sigma\tau} = \begin{cases} \frac{1}{2} s_\mu^{\sigma\tau} \text{ for } \rho = \sigma \text{ or } \tau, \\ 0 \text{ for } \rho \neq \sigma \text{ and } \tau, \end{cases} \\ \text{(c)} \quad s_\mu^{\rho\sigma} s_\nu^{\rho\sigma} = \begin{cases} e_\rho + e_\sigma \text{ for } \mu = \nu, \\ 0 \text{ for } \mu \neq \nu, \end{cases} \\ \text{(d)} \quad s_\mu^{\rho\sigma} s_\nu^{\sigma\tau} = \sum_{\lambda=1}^{\chi_{\rho\tau}} A_{\nu\lambda\mu}^{\rho\sigma\tau} s_\lambda^{\rho\sigma\tau} \text{ for } \rho \neq \tau, \\ \text{(e)} \quad s_\mu^{\rho\sigma} s_\nu^{\tau\nu} = 0 \text{ for } \rho, \sigma \neq \tau, \nu. \end{array} \right.$$

The proof was given above.

7. The algebra A is completely described by (22) so that the only data we need are the integers γ and $\chi_{\rho\sigma}$ and the real numbers $A_{\nu\lambda\mu}^{\rho\sigma\tau}$. Our next problem is to deduce from the r -number character of A some conditions which will limit these numbers more closely.

THEOREM 15. *If ρ, σ, τ , and ν are four distinct numbers then, under the above assumptions,*

$$(23) \quad A_{\nu\lambda\mu}^{\rho\sigma\tau} = A_{\lambda\mu\nu}^{\sigma\tau\rho} = A_{\mu\nu\lambda}^{\tau\rho\sigma} = A_{\lambda\nu\mu}^{\sigma\tau\rho} = A_{\nu\mu\lambda}^{\rho\tau\sigma} = A_{\mu\lambda\nu}^{\tau\sigma\rho},$$

$$(24) \quad \sum_{\nu=1}^{\chi_{\sigma\tau}} (A_{\nu\lambda\mu}^{\rho\sigma\tau} A_{\nu\lambda'\mu'}^{\rho\sigma\tau} + A_{\nu\lambda\mu'}^{\rho\sigma\tau} A_{\nu\lambda'\mu}^{\rho\sigma\tau}) = \frac{1}{2} \delta_{\lambda\lambda'} \delta_{\mu\mu'} \quad (\text{see footnote 10}),$$

$$(25) \quad \sum_{\mu=1}^{\chi_{\sigma\nu}} A_{\mu\theta_1}^{\rho\sigma\nu} A_{\kappa\mu\lambda}^{\sigma\tau\nu} = \sum_{\nu=1}^{\chi_{\rho\tau}} A_{\kappa\theta\nu}^{\rho\tau\nu} A_{\lambda\nu_1}^{\rho\sigma\tau}.$$

Proof: We substitute for a, b, c , and u in (4) the following combinations of the elements of the basis in Theorem 14: $s_\mu^{\rho\sigma}, s_\nu^{\sigma\tau}, e_\tau, s_\lambda^{\rho\tau}; s_\mu^{\rho\sigma}, s_\mu^{\sigma\tau}, s_\lambda^{\rho\tau}, e_\tau, s_\mu^{\rho\sigma}, s_\lambda^{\sigma\tau}, s_\kappa^{\tau\nu}, e_\rho$. The first substitution gives $A_{\nu\lambda\mu}^{\rho\sigma\tau} = A_{\mu\nu\lambda}^{\tau\rho\sigma}$ and (23) follows by cyclic permutation of the indices. (24) follows from the second substitution with the aid of (23). The third substitution leads to (25).

Although the relations (23), (24), and (25) are merely necessary and not sufficient, yet they will enable us in what follows, particularly in Part III, to determine completely all r -number algebras.

We now draw a few conclusions.

THEOREM 16. *If $\chi_{\rho\tau} \neq 0$, then under the above assumptions $\chi_{\rho\sigma} = \chi_{\sigma\tau}$.*

Proof: Since $\chi_{\rho\tau} \geq 1$, we can set $\lambda = \lambda' = 1$ in (24) and it follows that $2 \sum_{\nu=1}^{\chi_{\sigma\tau}} A_{\nu 1 \mu}^{\rho\sigma\tau} A_{\nu 1 \mu'}^{\rho\sigma\tau} = \frac{1}{2} \delta_{\mu\mu'}$. Therefore the vectors $2A_{\nu 1 \mu}^{\rho\sigma\tau}, \dots, 2A_{\chi_{\sigma\tau} 1 \mu}^{\rho\sigma\tau}$, where $\mu = 1, \dots, \chi_{\rho\sigma}$, form a normalized orthogonal system. Since their dimension

is $\chi_{\sigma\tau}$ and since there are $\chi_{\rho\sigma}$ of these vectors, $\chi_{\rho\sigma} \leq \chi_{\sigma\tau}$. Interchange of ρ and τ gives $\chi_{\sigma\tau} \leq \chi_{\rho\sigma}$, whence $\chi_{\rho\sigma} = \chi_{\sigma\tau}$.

THEOREM 17. *If A is irreducible then, under the above assumptions, all the $\chi_{\rho\sigma}$ are equal to each other and distinct from zero.*

Proof: Let N be the set of all ρ for which a chain $\rho_0 = 1, \rho_1, \dots, \rho_{\epsilon-1}, \rho_{\epsilon} = \rho$ ($\epsilon = 0, 1, 2, \dots$) exists such that every $\chi_{\rho_{\eta-1}\rho_{\eta}} \neq 0$ for $\eta = 1, \dots, \epsilon$. Let M be the set of all other ρ . N is not empty since $1 \in N$. If $\rho \in N$ and $\sigma \in M$, then $\chi_{\rho\sigma} = 0$. Hence $M^{\rho\sigma}$ contains only the element 0. $e = \sum_{\rho \in N} e_{\rho}$ is a unit element.

By (17) and (18) it follows that the condition $ex = \frac{1}{2}x$ is equivalent to the condition that, in the resolution (18) of x , only those addends $x_{\rho\sigma}$ occur for which $\rho \in N$ and $\sigma \in M$. Hence $x_{\rho\sigma} \in M^{\rho\sigma}$ and $x_{\rho\sigma} = 0$ so that $x = 0$. Therefore $N_1(e) = 0$ and A is the sum of the two orthogonal subalgebras $N_0(e)$ and $N_1(e)$. Since A is irreducible and $N_1(e) \neq 0$, then $N_0(e) = 0$. But every e_{ρ} , $\rho \in M$, is distinct from zero and belongs to $N_0(e)$ (since $ee_{\rho} = 0$). Hence M is empty and every ρ , $\rho = 1, \dots, \gamma$, belongs to N . Thus we can set up the chain $\rho_0 = 1, \rho_1, \dots, \rho_{\epsilon-1}, \rho_{\epsilon} = \rho$ for every ρ .

From Theorem 16 it follows that $\chi_{\rho_{\eta-1}\sigma} = \chi_{\rho_{\eta}\sigma}$ for every σ since $\chi_{\rho_{\eta-1}\rho_{\eta}} \neq 0$. Therefore $\chi_{\rho_0\sigma} = \chi_{\rho_{\epsilon}\sigma}$ and $\chi_{\rho\sigma} = \chi_{1\sigma}$. But we can write $\chi_{1\sigma}$ as $\chi_{\sigma 1}$. By the preceding argument it follows that $\chi_{1\sigma} = \chi_{11}$ so that $\chi_{\rho\sigma} = \chi_{11}$.

If we denote the common value of the $\chi_{\rho\sigma}$ by χ , then Theorem 14 provides the following condition on the dimension N of A :

$$(26) \quad N = \gamma + \frac{\gamma(\gamma - 1)}{2} \chi, \quad \chi > 0.$$

It should be noted that, while N is uniquely determined by A , this by no means needs to be the case for γ and χ as introduced in Theorems 12, 13, 14 and 17. We shall see later (Theorem 23 in Part II) that γ and χ are actually so determined.

Part II. Trace, Commutativity

8. The investigation of the concepts mentioned in the title of this part is not absolutely necessary for the logical connection of the entire paper since in Part III we shall determine directly from Part I all r -number algebras. And in the concrete instances taken up in Part III these concepts can be set up directly without reference to the general theory. But it is our desire to develop in general form for the sake of eventual generalizations our present concepts which are in themselves not uninteresting. Moreover we arrive at the proof of the determinateness of γ and χ . Finally we touch on the application of these concepts to the forming of quantum mechanical concepts.

All the following considerations arise from the setting up of the operation $O_a x = ax$, where $x \in A$. If A has a linear basis $a_1, \dots, a_N$ then a matrix $\|T_{\alpha\beta}(a)\|$, $\alpha, \beta = 1, \dots, N$, corresponds to O_a :

$$(27) \quad a a_\alpha = \sum_{\beta=1}^N T_{\alpha\beta}(a) a_\beta \quad (\beta = 1, \dots, N).$$

If the basis is altered then $\|T_{\alpha\beta}(a)\|$ is merely transformed as follows:

$$(28) \quad \begin{cases} a'_\alpha = \sum_{\beta=1}^N \psi_{\alpha\beta} a_\beta, \\ \|T'_{\alpha\beta}(a)\| = \|\psi_{\alpha\beta}\| \cdot \|T_{\alpha\beta}(a)\| \cdot \|\psi_{\alpha\beta}\|^{-1}. \end{cases} \quad \beta = 1, \dots, N,$$

Therefore the concepts which are invariant under transformations, such as trace $\|T_{\alpha\beta}(a)\| = \sum_{\alpha=1}^N T_{\alpha\alpha}(a)$ and the commutativity of $\|T_{\alpha\beta}(a)\|$ and $\|T_{\alpha\beta}(b)\|$ (that is, of O_a and O_b), are independent of the selection of the basis.

9. Because of the non-associativity of multiplication in A it need not be the case that $O_a O_b$ is either O_{ab} or O_{ba} , for otherwise it would be necessary that the relation $a(bx) = (ab)x$, that is, $[a, b, x] = 0$, hold for all x . From $ab = ba$ it does not follow that $O_a O_b = O_b O_a$ (that is, the commutativity of the corresponding matrices), for this would imply the relations $a(bx) = b(ax)$ and $a(xb) = (ax)b$, that is, $[a, x, b] = 0$ for all x . Consequently it is impossible to define even the general commutativity of multiplication as a concept of interchangeability in A .

Accordingly we call a and b *interchangeable* if the operations O_a and O_b are interchangeable, that is, if the matrices $\|T_{\alpha\beta}(a)\|$ and $\|T_{\alpha\beta}(b)\|$ are interchangeable, and hence if

$$(29) \quad [a, x, b] = 0$$

holds for all x . The *centrum* of A consists of those elements a with which all elements b are interchangeable and hence which are such that

$$(30) \quad [a, x, y] = 0$$

holds for all x and y . In this event it follows by (2) and (3) that

$$(31) \quad [x, a, y] = 0 \text{ and } [x, y, a] = 0.$$

With regard to the preceding definition of interchangeability it must be kept in mind that it is essential that the basis of A be finite. In important instances of quantum mechanics analogous to the r -number algebras the situation is quite otherwise. There the elements $a, b, c, \dots$ are infinite Hermitian matrices $A, B, C, \dots$ and our product

$$AB = \frac{A \times B + B \times A}{2},$$

where $A \times B$ is the usual associative non-commutative matrix multiplication. (See footnote 1 as well as the discussion in the introduction.) A simple calculation shows that

$$[A, B, C] = \frac{B \times (A \times C - C \times A) - (A \times C - C \times A) \times B}{4},$$

so that interchangeability occurs in our sense only when $A \times C - C \times A = \lambda 1$. It follows from this that $A \times C - C \times A = 0$ only for finite matrices but not for infinite matrices.¹² In generalizations to infinite r -number algebras it is necessary to exercise caution with regard to interchangeability. It is clear that the concept of centrum must remain unaltered.

THEOREM 18. *The elements b which are interchangeable with a form an algebra $C(a)$ which contains 1. If one represents a according to Theorem 6 and formula (13), then the algebras $N_1(e_1), \dots, N_1(e_\mu)$ are pairwise orthogonal and $C(a)$ is their (direct) sum.*

Proof: For the resolution (18) of a with respect to the elements $e_1, \dots, e_\mu$ it follows that $N_1(e_\nu) = M''$ and $N_1(e_{\nu'}) = M'''$. By Theorem 9 these two algebras (cf. Theorem 4) are orthogonal for $\nu \neq \nu'$. Consequently $N_1(e_1) + \dots + N_1(e_\mu)$ is an algebra which contains $1 = e_1 + \dots + e_\mu$.

To show that a is interchangeable with each element of $N_1(e_1) + \dots + N_1(e_\mu)$ it is sufficient to show that it is interchangeable with each element of every $N_1(e_{\nu'})$. Since $a = \xi_1 e_1 + \dots + \xi_\mu e_\mu$ it is sufficient to show that each $x \in N_1(e_{\nu'})$ is interchangeable with each e_ν . By the above argument it follows that $x \in N_1(e_\nu)$ for $\nu = \nu'$ and that $x \in N_0(e_\nu)$ for $\nu \neq \nu'$. We now have to show that $[x, y, e_\nu] = 0$ for all y , but by Theorem 3 it suffices to consider the cases $y \in N_0(e_\nu)$, $y \in N_1(e_\nu)$, and $y \in N_1(e_{\nu'})$. In other words it remains to show that $(xy)e = x(ye)$ for all six combinations of $x \in N_0(e)$ and $x \in N_1(e)$ with $y \in N_0(e)$, $y \in N_1(e)$, and $y \in N_1(e)$. But this follows immediately from Theorems 4 and 7.

Conversely, suppose that b is interchangeable with a . We set $x = e_{\nu'}$ in $[a, x, b] = 0$ and resolve b according to (18) with respect to the $e_1, \dots, e_\mu$ corresponding to a . By (17) one concludes immediately that $[e_\nu, e_{\nu'}, b^{\rho\sigma}]$, i.e., $(e_\nu e_{\nu'}) b^{\rho\sigma} - e_\nu (e_{\nu'} b^{\rho\sigma})$ differs from zero only if $\rho \neq \sigma$, $\nu = \rho$, and $\nu' = \sigma$ (or, what amounts to the same thing, $\nu = \sigma$ and $\nu' = \rho$) in which case it equals $-\frac{1}{4} b^{\rho\sigma}$. Consequently

$$[a, e_{\nu'}, b] = -\frac{1}{4} \sum_{\nu \neq \nu'} \xi_\nu b^{\nu\nu'}.$$

Therefore $b^{\nu\nu'} = 0$ for $\nu \neq \nu'$ if ξ_ν is not zero. So $b^{\nu\nu'} \neq 0$, $\nu \neq \nu'$ implies $\xi_\nu = 0$. By symmetry, it implies $\xi_{\nu'} = 0$. But this contradicts the fact that all the ξ_ν are distinct (note the discussion preceding Theorem 6). Hence

$$b = \sum_{\nu=1}^{\mu} b'' \text{ and } b \in (N_1(e_1) + \dots + N_1(e_\mu)).$$

¹² For example, A and C may be the matrices arising from the well-known quantum mechanical operators x and $i \frac{d}{dx}$.

THEOREM 19. *The centrum Z is an algebra which contains 1. Its degree is the number k of the irreducible components of A . (Note the end of paragraph 5.)*

Proof: By Theorem 18, the first assertion follows from the fact that Z is the common part of all $C(a)$. It can also be deduced directly from the identity¹³

$$[ab, x, y] = [a, bx, y] - [a, b, xy] + a[b, x, y] + [a, b, x]y.$$

If $A = B_1 + \cdots + B_k$ is the representation of A as the direct sum of irreducible algebras, then let $\bar{e}_1, \cdots, \bar{e}_k$ be the moduli of $B_1, \cdots, B_k$. $\bar{e}^2 = \bar{e}_i \neq 0$ and $\bar{e}_i \bar{e}_j = 0$ for $i \neq j$ (B_i and B_j orthogonal); if $x_i \in B_i$, then $\bar{e}_i x_i = x_i$ for $i = j$ and $\bar{e}_i x_j = 0$ for $i \neq j$ (also because of orthogonality). Thus $(\bar{e}_1 + \cdots + \bar{e}_k)x = x$ for $x = x_j$ and accordingly for every x . Hence $\bar{e}_1 + \cdots + \bar{e}_k = 1$. According to the discussion at the end of paragraph 5 it follows that $N_1(\bar{e}_i) = 0$ so that $A = N_0(\bar{e}_i) + N_1(\bar{e}_i) = N_1(1 - \bar{e}_i) + N_1(\bar{e}_i)$. Thus the resolution of $\bar{e}_i$ according to (13) is $1 \cdot \bar{e}_i + 0 \cdot (1 - \bar{e}_i)$ so that, by Theorem 18, $C(\bar{e}_i) = A$. Hence $\bar{e}_i$ is interchangeable with every b and $\bar{e}_i \in Z$. But $\bar{e}_1, \cdots, \bar{e}_k$ are linearly independent and the dimension of Z is greater than or equal to k .

Conversely, suppose that a belongs to Z . Then every polynomial in a belongs to Z and, in particular (because of (12)), the elements $e_1, \cdots, e_\mu$ of the resolution (13). We shall show that $e_1, \cdots, e_\mu$ are linear aggregates of $\bar{e}_1, \cdots, \bar{e}_k$ and hence that a is such an aggregate. It follows from this that the dimension of Z is k .

Let $e \in Z$ be a unit element. Its resolution according to (13) is $1 \cdot e + 0 \cdot (1 - e)$. Thus $C(e) = N_0(e) + N_1(e)$. Since $e \in Z$, $C(e) = A$, i.e., $N_0(e) + N_1(e) = A$. N_0 and N_1 are of course orthogonal. If we resolve the algebras N_0 and N_1 into their irreducible components: $N_0 = B'_1 + \cdots + B'_\alpha$, $N_1 = B''_1 + \cdots + B''_\beta$, then we have the resolution $A = B'_1 + \cdots + B'_\alpha + B''_1 + \cdots + B''_\beta$ which must be a permutation of $A = B_1 + \cdots + B_k$. But the modulus e of N_1 is the sum of the moduli of $B'_1, \cdots, B'_\alpha$ (cf. the corresponding situation relative to A in the preceding theorem). Hence e is the sum of the moduli of some of the B_i , that is, of some of the $\bar{e}_i$. Hence e is a linear aggregate of $\bar{e}_1, \cdots, \bar{e}_k$ as asserted above.

By Theorem 19, A is irreducible (that is, $k = 1$) when and only when the dimension of Z is unity, that is, when A consists of the elements $\lambda 1$ alone.

10. The trace of a is a (real) number defined by the equation

$$\text{trace } a = \Delta \text{ trace } \|T_{\alpha\beta}(a)\| = \Delta \sum_{\alpha=1}^N T_{\alpha\alpha}(a).$$

As we know, the trace is independent of the selection of the basis of A (cf. the end of paragraph 8). The normalization factor Δ is some positive number which we shall select later.

¹³ Cf. M. Zorn, Hamb. Abh., vol. 8, 1930, p. 123. This follows from the distributive law of multiplication alone.

It is clear that the trace is linear in a . The relation (which we do not prove)

$$O_{[a, b, c]} = O_b(O_a O_c - O_c O_a) - (O_a O_c - O_c O_a) O_b$$

implies the same relation for the corresponding matrices. From this it follows that $\text{trace } [a, b, c] = 0$ (since the trace of every matrix commutator vanishes). This is noteworthy because of its analogy with certain formulas of the quantum mechanical matrix-theory, but it will not be used in the sequel. We now consider other properties of the trace which are essential for the development of the quantum mechanical calculus.

THEOREM 20. *If e and e' are two unresolvable unit elements, there exists a resolution of 1 in pairwise orthogonal unresolvable unit elements: $e_1 + \dots + e_\gamma = 1$, where $e_1 = e$ and $\gamma \geq 2$, so that*

$$(32) \quad \begin{cases} e' = \alpha e_1 + (1 - \alpha) e_2 + x \\ x^2 = \alpha(1 - \alpha)(e_1 + e_2) \end{cases} \quad \begin{array}{l} (0 \leq \alpha \leq 1, x \in M^{12}), \\ (\text{Cf. Theorem 12}). \end{array}$$

Proof: We make the resolution $e' = x_0 + x_{\frac{1}{2}} + x_1$, where $x_0 \in N_0(e)$, $x_{\frac{1}{2}} \in N_{\frac{1}{2}}(e)$, and $x_1 \in N_1(e)$. Since e is unresolvable, it follows by Theorem 11 that $x_1 = \alpha e$. We resolve x_0 in the algebra $N_0(e)$ according to (13): $x_0 = \xi_2 e_2 + \dots + \xi_\gamma e_\gamma$. (We begin with the index 2 for formal reasons.) Since we can represent each e_ν as a sum of pairwise orthogonal unit elements (this sum being an element of $N(e_\nu)$ so that the unit elements for distinct ν are mutually orthogonal), we can arrive at the same representation with unresolvable unit elements. So we may assume that the preceding e_ν are unresolvable, but then the ξ_ν need no longer be distinct. The elements $e_2, \dots, e_\gamma$ are all orthogonal to e_1 and their sum is the modulus $1 - e$ of $N_0(e) = N_1(1 - e)$. Hence if $e_1 = e$ then $e_1 + \dots + e_\gamma = 1$ is a resolution of 1 of the desired sort. Since $N_{\frac{1}{2}}(e) = N_{\frac{1}{2}}(e_1) = M^{12} + \dots + M^{1\gamma}$ (cf. the analogous argument in the proof of Theorem 9) we have

$$e' = \alpha e_1 + \xi_2 e_2 + \dots + \xi_\gamma e_\gamma + x_2 + \dots + x_\gamma,$$

where $x_2 \in M^{12}, \dots, x_\gamma \in M^{1\gamma}$. By Theorem 12 we have that $x_\nu^2 = \eta_\nu(e_1 + e_\nu)$ and therefore that

$$\begin{aligned} e'^2 = & (\alpha^2 + \eta_2 + \dots + \eta_\gamma) e_1 + (\xi_2^2 + \eta_2) e_2 + \dots + (\xi_\gamma + \eta_\gamma) e_\gamma + (\alpha + \xi_2) x_2 \\ & + \dots + (\alpha + \xi_\gamma) x_\gamma + \sum_{\nu < \nu'} x_\nu x_{\nu'} \quad (\nu, \nu' = 2, \dots, \gamma). \end{aligned}$$

As was the case with e' , the terms of the summation belong to $M^{\nu\nu'}$ ($\nu, \nu' = 2, \dots, \gamma; \nu < \nu'$), while the other terms in this equation belong respectively to $M^{11}, M^{22}, \dots, M^{\gamma\gamma}, M^{12}, \dots, M^{1\gamma}$. Since $e'^2 = e'$, therefore $x_\nu x_{\nu'} = 0$. With the help of formulas (22 c, d), (23) and (24) one may readily calculate that $x_\nu(x_\nu x_{\nu'}) = \frac{1}{2} \eta_\nu x_{\nu'}$ (represent x_ν and $x_{\nu'}$ as linear aggregates of the bases $s_i^{\nu\nu'}$ and $s_i^{1\nu'}$). Hence either $x_{\nu'} = 0$ or $\eta_\nu = 0$ so that $x_\nu^2 = 0$ and $x_\nu = 0$. Hence at most one x_ν is different from zero and, by a permutation of the elements $e_2, \dots, e_\gamma$, we

may take it to be x_2 . Consequently $\eta_3 = \cdots = \eta_\gamma = 0$, $x_3 = \cdots = x_\gamma = 0$, and, since $e'^2 = e'$,

$$(33) \quad \begin{cases} \alpha^2 + \eta_2 = \alpha, \xi_2^2 + \eta_2 = \xi_2, \xi_3^2 = \xi_3, \dots, \xi_\gamma^2 = \xi_\gamma, \\ \alpha + \xi_2 = 1 \text{ for } x_2 \neq 0. \end{cases}$$

By (33) every ξ_ν is either zero or unity for $\nu = 3, \dots, \gamma$. A solution remains a solution if a ξ_ν with value unity is given the value zero instead. If $\xi_\nu = 1$, then e_ν and $e' - e_\nu$ are two orthogonal unit elements provided that $e' - e_\nu \neq 0$. Since e' is unresolvable, it must be the case that $e' = e_\nu$. By a permutation of the elements $e_2, e_3, \dots, e_\gamma$ it can be arranged that $e' = e_2$. This is in agreement with (32) when $\alpha = 0$ and $x = 0$. If, on the contrary, every ξ_ν is zero for $\nu = 3, \dots, \gamma$, then we may draw the following conclusion from (33): If $x = 0$, then e' has one of the values $e_1, e_2, e_1 + e_2$, or 0. But it is obvious that 0 cannot be a value of e' . Neither can $e_1 + e_2$ because of our hypothesis of unresolvability. Hence e is either e_1 or e_2 . This agrees with (32) in the case where α is 1 or 0 and $x = 0$. But if $x \neq 0$, then $\xi_2 = 1 - \alpha$, $\eta_2 = \alpha(1 - \alpha) > 0$ so that $0 < \alpha < 1$. This agrees with (32) in the case where α is neither 0 nor 1.

THEOREM 21. *If e and e' are non-orthogonal unit elements, $\text{trace}(ee') > 0$.*

Proof: By Theorem 10 we can represent e and e' as sums of unresolvable unit elements. Consequently it is sufficient to consider the case where e and e' are unresolvable. By (32) it follows that

$$ee' = \alpha e_1 + \frac{1}{2}x, \text{ where } x \in M^{12}, x^2 = \alpha(1 - \alpha).$$

If we now select the basis of A defined in Theorem 14, then it follows by (22) and the definition of trace that

$$\text{trace } e_\rho = \Delta \left(1 + \frac{1}{2} \sum_{\rho \neq \sigma} \chi_{\rho\sigma} \right) > 0, \quad \text{trace } s_\mu^{\rho\sigma} = 0 \quad (\rho, \sigma = 1, \dots, \gamma).$$

Hence if $\alpha > 0$, then $\text{trace}(ee') = \alpha/2 \text{ trace}(e_1) > 0$. But since $0 \leq \alpha \leq 1$, therefore the only alternative is $\alpha = 0$, which in turn implies $x^2 = 0$, $x = 0$, and therefore $ee' = 0$.

THEOREM 22. *If A is irreducible then Δ can be given such a positive value that $\text{trace}(e) = 1$ for every unresolvable unit element e .*

Proof: For every resolution $e_1 + \cdots + e_\gamma$ of 1 in pairwise orthogonal unresolvable unit elements it follows by the trace formulas in the proof of Theorem 21 and the remarks at the end of paragraph 7 that

$$\text{trace } e_\rho = \Delta \left(1 + \frac{(\gamma - 1)\chi}{2} \right) = \frac{\Delta N}{\gamma}, \quad \text{trace } s_\mu^{\rho\sigma} = 0.$$

Thus γ , χ , and all $\text{trace } e_\rho$ have the same values for every such resolution having the same $e_1 = e$. For any other unresolvable unit element e' it follows by (32) that $\text{trace } e' = \text{trace } e_1 = \text{trace } e$ since $\text{trace } x = 0$ as all $\text{trace } s_\mu^{12} = 0$. There-

fore every unresolvable unit element e has the same trace. Consequently γ and χ always have the same value. We may therefore set

$$(34) \quad \Delta = \frac{1}{1 + \frac{(\gamma - 1)\chi}{2}} = \frac{\gamma}{N} > 0$$

and arrive at the result¹⁴ that trace $e = 1$.

We shall always select Δ in this way if A is irreducible. We now formulate two immediate consequences.

THEOREM 23. *The numbers γ and χ defined at the end of paragraph 7 for an irreducible A are uniquely determined.*

Proof: Note the remark in the proof of Theorem 22.

THEOREM 24. *If A is irreducible a unit element e may be represented (in the sense of Theorem 10) as a sum of unresolvable unit elements. The number of addends is then always equal to trace e .*

Proof: The theorem follows immediately from Theorem 22.

The physical meaning of an r -number algebra can be determined from Theorems 6, 21, and 22. The elements $a, b, \dots$ correspond to the measurable quantities of a quantum system. The possible values of an element a arise from its resolution (13): $a = \xi_1 e_1 + \dots + \xi_\mu e_\mu$ and are the proper values $\xi_1, \dots, \xi_\mu$. If the value of a is measured and found to be ξ , then the state of the quantum system is described by the unit element e_ν . In the state e the expectation of an element b is $\frac{\text{trace}(eb)}{\text{trace } e}$. If $\eta_1 f_1 + \dots + \eta_j f_j$ is the resolution

of b by (13), then the probability that the measured value of b be η_i is $\frac{\text{trace}(ef_i)}{\text{trace } e}$. The case where e is unresolvable corresponds to a measurement which leads to a "reiner Fall."

Part III. Determination of all r -Number Algebras

11. We shall now seek all the solutions of the equations of condition (23), (24), and (25). These solutions will determine all r -number algebras which we shall assume throughout to be irreducible. We shall proceed stepwise in that we shall successively consider the cases $\gamma = 1, 2, 3$, and $\gamma \geq 4$.

The analysis starts with the linear basis of A described in Theorem 14 together with the multiplication rules (22). But before we actually begin, we remark that if the resolution theorem 14 is applicable to an r -number algebra then the condition $\chi_{\rho\sigma} = \chi > 0$ for irreducibility is not only necessary (as we already know) but also sufficient. Under this assumption every invariant subalgebra of A is either 0 or A . This means that if any $x \neq 0$ of A is given, all elements of A may be obtained from it by forming suitable linear aggregates of expressions

¹⁴ This conclusion is valid because of the previous unique determination of γ and χ .

of the form $((xa)b) \cdots k$, where $a, b, \cdots, k$ belong to A . For the formulas (18) and (19) show that in any case the elements $x_{\rho\sigma}$ can be so obtained. If some $x_{\rho\rho} \neq 0$ then we get an e_ρ , and if some $x_{\rho\sigma} \neq 0$ for $\rho \neq \sigma$ then we get an $x_{\rho\sigma}^2 \neq 0$. Therefore we can get an element $e_\rho + e_\sigma$ and from it we get $(e_\rho + e_\sigma)e_\rho = e_\rho$. Thus an e_ρ can be obtained in any case and from it we can get every element $e_\rho a_{\rho\tau} = \frac{1}{2} a_{\rho\tau}$ for $\rho \neq \tau$. Since we may select $a_{\rho\tau} \neq 0$ it follows that $a_{\rho\tau}^2 \neq 0$; this leads to $e_\rho + e_\tau$ and to $(e_\rho + e_\tau)e_\tau = e_\tau$. Hence we have all e_ρ and with them all $e_\rho \cdot 2s_\mu^{\rho\sigma} = s_\mu^{\rho\sigma}$ for $\rho \neq \sigma$. This completes the proof.

12. $\gamma = 1$. In this case $A = M^1$ and $e_1 = 1$ so that A consists merely of the elements $\lambda 1$. Since $\lambda 1 \pm \mu 1 = (\lambda \pm \mu)1$, $\lambda(\mu 1) = (\lambda\mu)1$, and $(\lambda 1)(\mu 1) = (\lambda\mu)1$, it follows that A is isomorphic with the algebra of real numbers.

13. $\gamma = 2$.¹⁵ In the basis of Theorem 14 we replace e_1 and e_2 by $e_1 + e_2 = 1$ and $e_1 - e_2 = s_{\chi+1}$, and we write $s_1, \cdots, s_\chi$ in place of $s_1^{12}, \cdots, s_\chi^{12}$. We observe that $N = \chi + 2$. This leads to the basis $1, s_1, \cdots, s_{N-1}$ of A and the multiplication rules

$$s_\mu^2 = 1, s_\mu s_\nu = 0 \text{ for } \mu \neq \nu.$$

Since $\chi > 0$ it follows that $N \geq 3$.

Conversely, A is an r -number algebra for each $N \geq 3$. This may be seen by verifying (I) and (III). In fact, if $a = \alpha 1 + \alpha_1 s_1 + \cdots + \alpha_{N-1} s_{N-1}$, then $a^2 = (\alpha^2 + \alpha_1^2 + \cdots + \alpha_{N-1}^2)1 + 2\alpha\alpha_1 s_1 + \cdots + 2\alpha\alpha_{N-1} s_{N-1}$. If $a \neq 0$ the coefficient of 1 is greater than zero, so that if $a^2 + b^2 + c^2 + \cdots = 0$ then $a = b = c = \cdots = 0$. But a^2 has the form $\beta 1 + \gamma a$ so that from the definition of 1 and the commutativity of multiplication it follows that $(a^2 b)a = a^2(ba)$.

14. $\gamma = 3$. It is now necessary to solve the equations (23) and (24). In this case (24) assumes the form

$$\sum_{\nu=1}^{\chi} (A_{\nu\lambda\mu}^{123} A_{\nu\lambda'\mu'}^{123} + A_{\nu\lambda\mu'}^{123} A_{\nu\lambda'\mu}^{123}) = \frac{1}{2} \delta_{\lambda\lambda'} \delta_{\mu\mu'}.$$

If we define χ matrices $B_1, \cdots, B_\chi$ by the equations

$$(35) \quad B_\lambda = \|B_{\lambda,\mu\nu}\|, \mu, \nu = 1, \dots, \chi; \quad B_{\lambda,\mu\nu} = 2A_{\nu\lambda\mu}^{123},$$

then¹⁶

$$(36) \quad B_\lambda^* B_{\lambda'} + B_{\lambda'}^* B_\lambda = 2\delta_{\lambda\lambda'} 1.$$

¹⁵ This case was first discussed by M. Zorn to whom the authors are indebted also for valuable assistance.

¹⁶ We denote the transposed matrix of B by B^* : $B = \|b_{\alpha\beta}\|$ and $B^* = \|b_{\beta\alpha}\|$. 1 is the unit matrix $1 = \|\delta_{\alpha\beta}\|$.

One can effect changes in the $A_{\nu\lambda\mu}^{123}$ (that is, the B_μ) by altering the bases of M^{22} , M^{13} , and M^{23} . These latter changes are unessential and, by Theorem 13 and (21), the transformations are

$$(37) \quad (s_\mu^{12})' = \sum_{\nu=1}^{\chi} \omega_{\mu\nu} s_\nu^{12}, \quad (s_\mu^{13})' = \sum_{\nu=1}^{\chi} \varphi_{\mu\nu} s_\nu^{13}, \quad (s_\mu^{23})' = \sum_{\nu=1}^{\chi} \psi_{\mu\nu} s_\nu^{23},$$

where the matrices $\|\omega_{\mu\nu}\|$, $\|\varphi_{\mu\nu}\|$, and $\|\psi_{\mu\nu}\|$ are orthogonal. By (22d) the elements $A_{\nu\lambda\mu}^{123}$ transform according to the relation

$$(38) \quad (A_{\nu\lambda\mu}^{123})' = \sum_{\nu', \lambda', \mu'=1}^{\chi} \omega_{\mu\mu'} \varphi_{\lambda\lambda'} \psi_{\nu\nu'} A_{\nu'\lambda'\mu'}^{123},$$

or, in terms of the B_μ ,

$$(39) \quad B'_\lambda = \sum_{\lambda'=1}^{\chi} \varphi_{\lambda\lambda'} U^* B_{\lambda'} W,$$

where $U = \|\omega_{\mu\nu}\|$, $V = \|\varphi_{\mu\nu}\|$, and $W = \|\psi_{\mu\nu}\|$ are unitary matrices. One sees immediately that for every solution of (36) there is another resulting from (39). We shall regard such solutions as being not essentially distinct.

The solution of (36) was obtained in a well-known paper¹⁷ by A. Hurwitz who showed in particular that χ must be 1, 2, 4, or 8. Since a part of Hurwitz's investigation (unimportant for his main purpose but important for our discussion in paragraph 15) contains an essential error,¹⁸ it is convenient to carry out the entire argument independently.

If in (36) $\lambda = \lambda'$, then $B_\lambda^* B_\lambda = 1$ so that all B_λ are orthogonal. If we set $C_\lambda = B_\chi^* B_\lambda$, $\lambda = 1, \dots, \chi - 1$, then the C_λ are also all orthogonal. If $\lambda \neq \lambda'$ then (36) may be written in the form $B_\lambda^* B_{\lambda'} = -B_{\lambda'}^* B_\lambda$ or $B_\lambda^* B_\lambda B_{\lambda'}^* B_{\lambda'} = -1$. Hence if $\lambda' = \chi$ ($\lambda \neq \lambda' = \chi$) then $C_\lambda^2 = -1$ and if $\lambda, \lambda' \neq \chi$ ($\lambda \neq \lambda'$) then $(C_\lambda^* C_{\lambda'})^2 = -1$ since $C_\lambda^* C_{\lambda'} = B_\lambda^* B_{\lambda'}$. Or again, $C_\lambda^* \cdot C_\lambda C_{\lambda'}^* \cdot C_{\lambda'} = -1$ while $C_\lambda^* \cdot C_{\lambda'}^* C_\lambda \cdot C_{\lambda'} = C_{\lambda'}^2 \cdot C_\lambda^2 = 1$ so that $C_\lambda C_{\lambda'}^* = -C_{\lambda'}^* C_\lambda$, or (after multiplication on both sides by $C_{\lambda'}$) $C_\lambda C_{\lambda'} = -C_{\lambda'} C_\lambda$. Hence the C_λ 's are orthogonal,

$$(40) \quad C_\lambda^2 = -1, C_\lambda C_{\lambda'} = -C_{\lambda'} C_\lambda \text{ for } \lambda \neq \lambda',$$

where $\lambda, \lambda' = 1, \dots, \chi - 1$.

If B is any orthogonal matrix, then $B_1 = BC_1, \dots, B_{\chi-1} = BC_{\chi-1}$,

¹⁷ "Über die Composition der quadratischen Formen von beliebig vielen Variablen," Goett. Nachr., 1898, p. 309.

¹⁸ Namely, the assertion that for a given χ all solutions B'_λ may be obtained from one particular B_λ by means of the relation $B'_\lambda = U^* B_\lambda W$ (that is, without the use of φ in (39)), this relation assuming the form $C'_\lambda = W^* C_\lambda W$ in the case where $C_\lambda = B_\chi^* B_\lambda$. (Cf. below.) One sees that this is incorrect in the following way: in case χ is equal to 4 or 8, then $C_1 \cdot \dots \cdot C_{\chi-1} = \pm 1$. This equation is invariant under the above transformation. Yet the left side has the opposite sign in the solution $B'_1 = -B$, $B'_2 = B$, $\dots$, $B'_\chi = B_\chi$. (Cf. footnote 17, pp. 313-315; our B, C, U, W are there denoted by A, B, P, Q .)

$B_x = B$ is the general solution of (36) in case $C_1, \dots, C_{x-1}$ is the general solution of (40).

15. Let G_x be that group whose modulus may be denoted by $\bar{1}$ and which possesses χ generators $-\bar{1}, \bar{\Gamma}_1, \dots, \bar{\Gamma}_{x-1}$ which satisfy the following relations:

$$(41) \quad \begin{cases} (-\bar{1})^2 = \bar{1}, & (-\bar{1})\bar{\Gamma}_\lambda = \bar{\Gamma}_\lambda(-\bar{1}), \\ \bar{\Gamma}_\lambda^2 = -1, & \bar{\Gamma}_\lambda\bar{\Gamma}_{\lambda'} = (-\bar{1})\bar{\Gamma}_{\lambda'}\bar{\Gamma}_\lambda \text{ for } \lambda \neq \lambda'. \end{cases}$$

It is obvious that all the elements of this group may be written as $\bar{\Gamma}_{\lambda_1} \dots \bar{\Gamma}_{\lambda_r}$ or as $(-\bar{1})\bar{\Gamma}_{\lambda_1} \dots \bar{\Gamma}_{\lambda_r}$, $r = 0, \dots, \chi - 1$, $1 \leq \lambda_1 < \dots < \lambda_r \leq \chi - 1$. Hence G_x is finite and has 2^x elements. The problem of solving (40) reduces to that of finding¹⁹ all representations of G_x by χ -dimensional orthogonal matrices in which $-\bar{1}$ is represented by the matrix -1 . $\bar{\Gamma}_\lambda$ is then represented by C_λ , $\lambda = 1, \dots, \chi - 1$.

We next determine the "irreducible" representations of G_x (cf. footnote 19 for all concepts and theorems in the theory of representations), and, of course, merely the complex-unitary ones. Since no C_λ exist for $\chi = 1$ (the solution is simply that B_1 must be orthogonal) we assume that $\chi \geq 2$.

The order of G_x is 2^x . The commutators of G_x are $\bar{1}$ and $-\bar{1}$ and they form an invariant subgroup of order 2 so that the index of G_x is 2^{x-1} . An element of the group can be transformed only into itself and its product with -1 , the latter being impossible for $\bar{1}$ and $-\bar{1}$ as well as for $\bar{\Gamma}_1 \dots \bar{\Gamma}_{x-1}$, $(-\bar{1})\bar{\Gamma}_1 \dots \bar{\Gamma}_{x-1}$ if χ is even. Hence 2 or 4 equivalence classes each contain one element according as χ is odd or even, while all others contain two elements. Their number is therefore $2^{x-1} + 1$ or $2^{x-1} + 2$. (The reader may complete the simple discussion of G_x .) Hence G_x has $2^{x-1} + 1$ or $2^{x-1} + 2$ inequivalent irreducible representations of which 2^{x-1} are of order 1. These are the commutative representations. But since every commutator is $\bar{1}$ or $-\bar{1}$, and is therefore representable by the matrices 1 or -1 , these commutative representations are exactly those in which $-\bar{1}$ is represented by 1. The representations which are of interest to us are therefore those of order $\neq 1$. From the above discussion it follows that there are one or two such representations.

The orders of all irreducible representations are divisors of 2^x and the sum of the squares of these orders is 2^x , of which the representations of order 1 contribute $2^{x-1} \cdot 1^2 = 2^{x-1}$. The sum of the squares of the orders of the remaining representations is $2^x - 2^{x-1} = 2^{x-1}$. If χ is odd there is only one such representation and its order is $2^{\frac{x-1}{2}}$. If χ is even there are two and their orders are divisors of 2^x . Hence their orders are powers of 2 and the sum of the squares of their orders can be 2^{x-1} only if each of them is $2^{\frac{x-2}{2}}$.

¹⁹ The G. Frobenius—I. Schur "theory of representations," which is prominent in what follows, may be found in I. Schur's treatise "Neue Begründung der Theorie der Gruppencharaktere," Berliner Ber., 1905, p. 406.

The (real) representation we are seeking must be made up of these irreducible (complex) representations and its order χ must be some multiple of the common degrees, i.e., $2^{\frac{\chi-1}{2}}$ or $2^{\frac{\chi-2}{2}}$. It follows immediately that $\chi = 2, 4$, or 8 . All these values are even, $\chi/2^{\frac{\chi-2}{2}} = 2, 2$, or 1 , and our representation consists of $2, 2$, or 1 irreducible parts. But since we wish to form a real representation from complex representations, we must keep the reality conditions in mind.²⁰

A complex irreducible representation can have in this regard one of three properties: 1) it may be equivalent to a real representation, 2) it may be equivalent to no real representation but instead be equivalent to its own complex-conjugate representation, or 3) it may be inequivalent to its complex-conjugate representation. In order to distinguish these cases one must form the sum of the traces (characteristics) of the matrices representing the squares of all the elements of the group. According as this sum is greater than, equal to, or less than 0 the situation is described by 1), 3), or 2).

In G_χ the square of every element is $\bar{1}$ or $-\bar{1}$. Suppose α of the elements have $\bar{1}$ for their square and suppose β of the elements have $-\bar{1}$ for their square. Then $\alpha + \beta = 2\chi$. The matrix representations of these squared elements are 1 or -1 and their trace is $2^{\frac{\chi-2}{2}}$ or $-2^{\frac{\chi-2}{2}}$. The sum in question is therefore $2^{\frac{\chi-2}{2}}(\alpha - \beta)$ and our criterion reduces merely to the questions whether α is greater than, equal to, or less than β . Enumeration shows that the relations $=$, $<$, and $>$ hold for the values $2, 4$, and 8 respectively of χ . Hence the cases 3), 2), and 1) exist. For $\chi = 2$ each of our two irreducible representations is inequivalent to its own complex-conjugate. Hence each is equivalent to the other so that we may describe the two irreducible representations as complex-conjugate to one another. The real representation we are seeking must therefore contain each of them exactly once. For $\chi = 4$ each irreducible representation is in case 2) and the real representation sought must contain each of them an even number of times. Hence it contains one of them twice but does not contain the other. For $\chi = 8$ each irreducible representation can be described as real so that both may be used as solutions. We may summarize our discussion by saying that for $\chi = 2$ there is one solution, and for $\chi = 4$ or 8 there are two inequivalent solutions.

Since $C_1 \cdots C_{\chi-1}$ is interchangeable with all C_λ (since χ is even), it must be of the form $\xi \cdot 1$ in each irreducible representation, and since its square is 1 , it must be ± 1 . Since the solution for $\chi = 4$ consists of two equivalent irreducible matrices and since the solution for $\chi = 8$ is irreducible, we have for $\chi = 4$ or 8 that

$$(42) \quad C_1 \cdots C_{\chi-1} = \pm 1.$$

²⁰ Cf. G. Frobenius and I. Schur, Berliner Ber., 1906, p. 186.

Since an equivalence transformation does not alter this relation, while the transformation $C'_1 = -C_1$, $C'_2 = C_2$, $\dots$, $C'_{\chi-1} = C_{\chi-1}$ interchanges the $+$ and $-$ signs in it, it follows that this transformation carries one solution over into the other. (Cf. footnote 18.)

If, for $\chi = 2$, we mean by s_λ simply $s_\lambda = 1$, and for $\chi = 4$ or 8 we mean the same or $s_1 = -1$, $s_2 = \dots = s_\chi = 1$, then the most general solution C'_λ arises from a solution C_λ by the transformation $C'_\lambda = s_\lambda W^* C_\lambda W$, where W is orthogonal. Hence $B'_\lambda = s_\lambda B W^* C_\lambda W$ (the relation also holds for $\lambda = \chi$ if we set $C_\chi = 1$: $B'_\chi = s_\chi B W^* W = B$). With the help of (39) it follows that $U = W B^*$ and $\varphi_{\mu\nu} = s_\mu \delta_{\mu\nu}$.

To sum up: Solutions exist for $\chi = 1, 2, 4$, and 8 , but for each of these values there is essentially only one solution in the sense of (37), (38), and (39). In particular we can set U and V equal to 1 for $\chi = 1$, we can set $V = 1$ for $\chi = 2$, and we can set $V = 1$ or $V = \|\bar{s}_\mu \delta_{\mu\nu}\|$ ($\bar{s}_1 = -1$, $\bar{s}_2 = \dots = \bar{s}_\chi = 1$) for $\chi = 4$ or 8 .

16. Having shown that there is only one solution, we may better find it in another way. (36) may also be reformulated (and this was the original formulation by A. Hurwitz). In fact, (36) is equivalent to

$$(43) \quad \begin{cases} (x_1^2 + \dots + x_\chi^2)(y_1^2 + \dots + y_\chi^2) = z_1^2 + \dots + z_\chi^2, \\ z_\lambda = \sum_{\mu, \nu=1}^{\chi} B_{\lambda, \mu\nu} x_\mu y_\nu \end{cases} \quad (\lambda = 1, \dots, \chi)$$

($x_1, \dots, x_\chi$ and $y_1, \dots, y_\chi$ being two sets of independent variables). The solution may be found by the well-known substitution

$$(44) \quad \begin{cases} X = x_1 1 + x_2 i_2 + \dots + x_\chi i_\chi, & Y = y_1 1 + y_2 i_2 + \dots + y_\chi i_\chi, \\ Z = z_1 1 + z_2 i_2 + \dots + z_\chi i_\chi, \\ X Y = Z \end{cases}$$

in case $i_1 = 1$, $i_2, \dots, i_\chi$ constitute the basis of the following hypercomplex number systems: for $\chi = 1$, the real numbers; for $\chi = 2$, the ordinary complex numbers; for $\chi = 4$, the quaternions;²¹ for $\chi = 8$, the so-called Cayley numbers or quasi-quaternions.²²

²¹ Quaternions satisfy the relations $i_2^2 = i_3^2 = i_4^2 = -1$; $i_2 i_3 = -i_3 i_2 = i_4$; $i_3 i_4 = -i_4 i_3 = i_2$; $i_4 i_2 = -i_2 i_4 = i_3$.

²² These were introduced by Cayley, *Corresp. Math. Phys.*, vol. 1, 1843, p. 452. (Published by P. H. Fuss). A summary, and at the same time historically exhaustive, description of them has been given by L. E. Dickson, *Annals. of Math.*, vol. 20, 1918, p. 155 and particularly pp. 157-159. A detailed discussion of their algebraic properties has been given by M. Zorn, cf. footnote 13. Quasi-quaternions are defined by $i_2^2 = \dots = i_8^2 = -1$ and each of the systems $i_2, i_3, i_6; i_3, i_4, i_6; i_4, i_5, i_7; i_5, i_6, i_8; i_6, i_7, i_2; i_7, i_8, i_3; i_8, i_2, i_4$ (together with $i_1 = 1$) has the formal properties of quaternions. They form a non-associative system of quantities.

If we denote the complex-conjugate of such a hypercomplex number $X = x_1 + x_2 i_2 + \cdots + x_\chi i_\chi$ by $\bar{X} = x_1 - x_2 i_2 - \cdots - x_\chi i_\chi$ and its real part by $R(X) = x_1$, then $R(X i_\mu) = x_\mu$, $\mu = 1, \cdots, \chi$. Consequently, if we note that $R(XY) = R(YX)$,

$$(45) \quad 2A_{\lambda\mu}^{123} = B_{\lambda,\mu\nu} = R(i_\mu i_\nu \bar{i}_\lambda) = R(i_\nu \bar{i}_\lambda i_\mu).$$

By (45) it is possible to calculate these $A_{\lambda\mu}^{123}$ with the help of the multiplication tables of the above hypercomplex number systems, but we shall not go into this. The general solution $(A_{\lambda\mu}^{123})'$ arises from (45) by (38) where we can keep the restriction mentioned at the end of paragraph 15.

It follows from the rules (22) together with (43) or (45) that the algebra A is isomorphic with a system of matrices $A, B, C, \cdots$ of order 3, wherein the operations $a \pm b$, λa , and ab correspond to the matrix operations $A \pm B$, λA , and $\frac{A \times B + B \times A}{2}$ ($A \times B$ is the usual matrix multiplication). Again, e_ρ corresponds to that matrix which has the $\rho\rho$ -element 1 but otherwise consists of zeros, and $s_\mu^{\rho\sigma}$ ($\rho < \sigma$) corresponds to that matrix which has i_μ for its $\rho\sigma$ -element and $\bar{i}_\mu$ for its $\sigma\rho$ -element but otherwise consists of zeros. Consequently A is the algebra of all third order Hermitian matrices, $\|X_{\rho\sigma}\|$, where $\rho, \sigma = 1, 2, 3$ and $X_{\rho\sigma} = \bar{X}_{\sigma\rho}$ (the diagonal elements $X_{\rho\rho}$ being real numbers), with the above addition and multiplication rules. Hence the matrix elements for $\chi = 1, 2, 4$, or 8 respectively are real numbers, complex numbers, quaternions, and quasi-quaternions and $N = \gamma + \frac{\gamma(\gamma-1)}{2} \chi = 3 + 3\chi = 6, 9, 15, 27$.

17. Conversely, if A is one of the above algebras, we shall ascertain whether or not it is an r -number algebra. Since we can form the system of Hermitian matrices described in the last paragraph for every order (and not merely of order three as was done there), and since we shall need the more general matrices later on, we generalize the problem as follows: A consists of all Hermitian matrices of order γ subject to the above addition and multiplication rules and having real numbers, complex numbers, quaternions, and quasi-quaternions respectively as elements. We provisionally assume merely that $\gamma \geq 3$ (since $\gamma = 1$ or 2 leads merely to special cases of the results obtained in paragraphs 12 and 13).

It is obvious that it is necessary to verify only (I) and either (II) or (III).

(I) follows from the fact that if $A = \|X_{\rho\sigma}\|$, then $\text{trace}(A^2) = \sum_{\rho, \sigma=1}^{\gamma} |X_{\rho\sigma}|^2$ and hence is real and, for $A \neq 0$, positive. (II) is evident for the first three cases ($\chi = 1, 2$, and 4) since the matrix elements belong to associative number systems so that $A^2, A^3, \cdots$ result from our multiplication AB just as with the usual matrix multiplication $A \times B$ (cf. footnote 1). It remains to consider only the fourth case ($\chi = 8$). We shall soon see that there is no solution for $\gamma \geq 4$ so

that only the case $\gamma = 3$ is in question. Instead of proving (III) we could prove its generalization (4) (which leads to (III) if one replaces a, b, c, u by a, a, a, b) which has the advantage of being linear in all four variables a, b, c, u . It suffices for this purpose to substitute all combinations of the 27 elements of the linear basis. This would lead actually to $27^4 = 531,441$ combinations on the basis of the multiplication rules and footnote 22. Yet this number may be reduced by a few devices so that the investigation is possible without too many possibilities. It results that (4) holds in all cases. We omit the proof which may be found in a subsequent paper of A. A. Albert in this journal. Thus for $\gamma = 3$ the fourth case also leads to a solution.

18. $\gamma \geq 4$. Since $N_1(e_1 + e_2 + e_3) = \sum_{\rho, \sigma=1}^3 M^{\rho\sigma}$ is a subalgebra of A which has the same χ as A but for which $\gamma = 3$, we again have $\chi = 1, 2, 4$, or 8 . We shall examine these cases successively.

$\gamma \geq 4, \chi = 8$. If in (25) we set ρ, σ, τ, ν respectively equal to 1, 2, 3, and 4, and introduce the following matrices (as in (35)):

$$\begin{aligned}\tilde{B}_\lambda &= \|\tilde{B}_{\lambda, \mu\nu}\|, \tilde{B}_{\lambda, \mu\nu} = 2A_{\mu\nu\lambda}^{123}; \tilde{C}_\lambda = \|\tilde{C}_{\lambda, \mu\nu}\|, \tilde{C}_{\lambda, \mu\nu} = 2A_{\nu\mu\lambda}^{124}; \\ \tilde{D}_\lambda &= \|\tilde{D}_{\lambda, \mu\nu}\|, \tilde{D}_{\lambda, \mu\nu} = 2A_{\lambda\mu\nu}^{134}; \tilde{E}_\lambda = \|\tilde{E}_{\lambda, \mu\nu}\|, \tilde{E}_{\lambda, \mu\nu} = 2A_{\lambda\nu\mu}^{234}.\end{aligned}$$

then it follows from (25) that

$$(46) \quad \tilde{C}_i \tilde{E}_\kappa^* = \tilde{D}_\kappa \tilde{B}_i^*,$$

and all of these matrices are orthogonal. (The orthogonality follows from (24), with either $\lambda = \lambda'$ or $\mu = \mu'$, together with (23).) If now we replace i by i' in (46) and form the transposed matrices, we have

$$\tilde{E}_\kappa \tilde{C}_{i'}^* = \tilde{B}_{i'} \tilde{D}_\kappa^*.$$

Multiplication of (46) on the right by this last equation gives

$$\tilde{C}_i \tilde{C}_{i'}^* = \tilde{D}_\kappa \tilde{B}_i^* \tilde{B}_{i'} \tilde{D}_\kappa^*.$$

The left side of this equation is independent of κ so that

$$\tilde{D}_\kappa \tilde{B}_i^* \tilde{B}_{i'} \tilde{D}_\kappa^* = \tilde{D}_{\kappa'} \tilde{B}_i^* \tilde{B}_{i'} \tilde{D}_{\kappa'}^*, \quad \tilde{D}_\kappa^* \tilde{D}_{\kappa'} \tilde{B}_i^* \tilde{B}_{i'} = \tilde{B}_i^* \tilde{B}_{i'} \tilde{D}_\kappa^* \tilde{D}_{\kappa'}.$$

Hence every $\tilde{B}_i^* \tilde{B}_{i'}$ commutes with every $\tilde{D}_\kappa^* \tilde{D}_{\kappa'}$.

We now have, with regard to (23),

$$\tilde{D}_{\lambda, \mu\nu} = 2A_{\lambda\mu\nu}^{134} = 2A_{\lambda\nu\mu}^{413}.$$

Hence in the subalgebra $N_1(e_1 + e_3 + e_4) = \sum_{\rho, \sigma=1,3,4} M^{\rho\sigma}$ of A the matrix $\tilde{D}_\lambda = \|\tilde{D}_{\lambda, \mu\nu}\|$ plays the same rôle as did the matrix $B_\lambda = \|B_{\lambda, \mu\nu}\|$ in paragraphs 14 to 17 in the r -number algebra with $1 = e_1 + e_2 + e_3$. (The indices 1, 2, 3 are

replaced by 4, 1, 3.) The matrices $\tilde{D}_\chi^* \tilde{D}_\kappa$, $\kappa = 1, \dots, \chi - 1$, therefore play the rôle of the C_κ and therefore form, for our $\chi = 8$, an irreducible system (cf. paragraph 15). Since every $\tilde{B}_i^* \tilde{B}_{i'}$ is interchangeable with every $\tilde{D}_\chi^* \tilde{D}_\kappa$, they all have the form $\xi_{i, \kappa} 1$. But with regard to (23),

$$\tilde{B}_{\lambda, \mu\nu} = 2A_{\mu\nu\lambda}^{123} = 2A_{\lambda\mu\nu}^{231}.$$

Hence in the subalgebra $N_1(e_1 + e_2 + e_3) = \sum_{\rho, \sigma=1}^3 M^{\rho\sigma}$ of A the matrix $\tilde{B}_\lambda = \|\tilde{B}_{\lambda, \mu\nu}\|$ plays the same rôle as did the matrix $B_\lambda = \|B_{\lambda, \mu\nu}\|$ in paragraphs 14 to 17 in the r -number algebra with $1 = e_1 + e_2 + e_3$. (The indices 1, 2, 3 are replaced by 2, 3, 1.) The matrices $\tilde{B}_\chi^* \tilde{B}_\kappa$, $\kappa = 1, \dots, \chi - 1$, play the rôle of the C_κ and form, for our $\chi = 8$, an irreducible system. Therefore the form $\xi_{\chi\kappa} 1$ with the order $\chi = 8 > 1$ is impossible.

Thus in the above case there exists no solution.

19. $\gamma \geq 4$, $\chi = 4$. In (45) we obtained a solution for $2A_{\nu\lambda\mu}^{123}$ in the form $R(i_\nu \bar{i}_\lambda i_\mu)$, where i_1, i_2, i_3, i_4 compose the basis of the system of quaternions. (Cf. footnote 21.) By (38) the most general solution results from it if we replace

$$i_\nu \text{ by } \sum_{\nu'=1}^4 \psi_{\nu\nu'} i_{\nu'}, \quad i_\mu \text{ by } \sum_{\mu'=1}^4 \varphi_{\mu\mu'} i_{\mu'}, \quad i_\lambda \text{ by } \sum_{\lambda'=1}^4 \omega_{\lambda\lambda'} i_{\lambda'}$$

(we make no use of the restriction proved at the end of paragraph 15). But every orthogonal transformation $\sum_{\kappa'=1}^4 \sigma_{\kappa\kappa'} i_{\kappa'}$ of the basis of the quaternion system may also be represented in the form $X' i_\kappa \cdot X''$, where X' and X'' represent certain quaternions of absolute value unity and $\cdot$ denotes nothing at all or else the sign for the complex conjugate.²³ The most general expression for $2A_{\nu\lambda\mu}^{123}$ is therefore (since $R(XY) = R(YX)$)

$$2A_{\nu\lambda\mu}^{123} = R(Y^{(123)} i_\nu^{(312)} Y^{(231)} i_\lambda^{(123)} Y^{(312)} i_\mu^{(231)})$$

where the Y 's are certain quaternions of absolute value unity and where the $\cdot$'s are independent of one another. Since $R(XY) = R(YX)$ it follows by (23) that this equation also holds for $2A_{\nu\lambda\mu}^{231}$ and $2A_{\nu\lambda\mu}^{312}$ if the indices 1, 2, and 3 are correspondingly permuted on the right side. Hence we may set

$$(47) \quad 2A_{\nu\lambda\mu}^{\rho\sigma\tau} = R(Y^{(\rho\sigma\tau)} i_\nu^{(\tau\rho\sigma)} Y^{(\sigma\tau\rho)} i_\lambda^{(\rho\sigma\tau)} Y^{(\tau\rho\sigma)} i_\mu^{(\sigma\tau\rho)})$$

where the Y 's are certain quaternions of absolute value unity and where the $\cdot$'s are independent of one another. Consequently we shall limit ourselves to the

²³ We will write $\cdot = 0$ if $\cdot$ denotes nothing, and $\cdot = -$ if it denotes the complex conjugate. The representation of rotations by means of quaternions can be verified easily.

consideration of such values of ρ , σ , and τ which may be transformed by a cyclic permutation into a monotonically increasing sequence, and we shall call them allowable sequences. (We have treated the permutations of 1, 2, and 3 in the same manner. Since ρ , σ , τ or τ , σ , ρ is always allowed in this sense, it follows that the allowable $A_{\lambda\mu\nu}^{\rho\sigma\tau}$ are sufficient for the purpose of the multiplication rules in (22d).)

On the other hand the generality of (47) can be somewhat lessened in that we can subject every $M^{\rho\sigma}$ to an arbitrary transformation in the sense of Theorem 13, (21), and (37). This is exhibited in the formulas (47) for $2A_{\nu\lambda\mu}^{\rho\sigma\tau}$, $2A_{\nu\lambda\mu}^{\sigma\tau\rho}$, $2A_{\nu\lambda\mu}^{\tau\rho\sigma}$, or $2A_{\nu\lambda\mu}^{\sigma\rho\tau}$, $2A_{\nu\lambda\mu}^{\rho\tau\sigma}$, $2A_{\nu\lambda\mu}^{\tau\sigma\rho}$ as a transformation of the corresponding i_μ , i_λ , i_ν and hence as that of the i_κ with $(\sigma\tau\rho)$ or $(\rho\tau\sigma)$. If we transform all the $M^{\rho\sigma}$ simultaneously so that in $M^{\rho\sigma}$ i_κ goes over into $X_1^{(\rho\sigma)} i_\kappa^{(\rho\sigma)} X_2^{(\rho\sigma)}$, and if for the sake of simplicity we identify $\rho\sigma$ and $\sigma\rho$, then we have (since $\overline{XY} = \overline{YX}$)

$$(48) \quad \left\{ \begin{array}{l} (a) \quad i_\kappa^{(\rho\sigma\tau)} \text{ is replaced by } i_\kappa^{(\rho\sigma\tau) \cdot (\sigma\tau)}, \\ (b) \quad Y^{(\rho\sigma\tau)} \text{ is replaced by } X_I^{(\rho\sigma)} Y^{(\rho\sigma\tau)} X_{II}^{(\sigma\tau)}, \text{ where} \\ \quad X_I^{(\rho\sigma)} = X_2^{(\rho\sigma)} \text{ or } \bar{X}_I^{(\rho\sigma)} \text{ when } (\sigma\tau\rho) = 0 \text{ or } \overline{}, \\ \quad X_{II}^{(\sigma\tau)} = X_I^{(\sigma\tau)} \text{ or } \bar{X}_2^{(\sigma\tau)} \text{ when } (\tau\rho\sigma) = 0 \text{ or } \overline{}. \end{array} \right.$$

We now consider (25). (Since ρ , σ , τ ; ρ , σ , ν ; ρ , τ , ν ; and σ , τ , ν are to be allowable, it must be the case that every ρ , σ , τ , ν goes over into a monotonically increasing sequence by a cyclic permutation and hence must be allowable.) If we substitute

$$R(XY) = R(YX), \sum_{\kappa=1}^4 R(X i_\kappa^{(1)}) R(Y i_\kappa^{(2)}) = R(X \bar{Y}^{(1)} \bar{Y}^{(2)})$$

in (47), we get

$$\begin{aligned} & R([Y^{(\sigma\nu\rho)} i_\theta^{(\rho\sigma\nu)} Y^{(\nu\rho\sigma)} i_\tau^{(\sigma\nu\rho)} Y^{(\rho\sigma\nu)}] \cdot [Y^{(\nu\sigma\tau)} i_\lambda^{(\tau\nu\sigma)} Y^{(\sigma\tau\nu)} i_\kappa^{(\nu\sigma\tau)} Y^{(\tau\nu\sigma)}] \overline{(Y^{(\rho\sigma\tau)} i_\nu^{(\tau\rho\sigma)} Y^{(\sigma\tau\rho)})}) \\ &= R([Y^{(\rho\tau\nu)} i_\kappa^{(\nu\rho\tau)} Y^{(\tau\nu\rho)} i_\theta^{(\rho\tau\nu)} Y^{(\nu\rho\tau)}] \cdot [Y^{(\tau\rho\sigma)} i_\tau^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)} i_\lambda^{(\tau\rho\sigma)} Y^{(\sigma\tau\rho)}] \overline{(Y^{(\tau\rho\sigma)} i_\nu^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)})}). \end{aligned}$$

Since we are forming linear aggregates we can set a free quaternion variable Z_θ in place of i_θ and likewise free quaternion variables Z_i , Z_λ , and Z_κ in place of i_i , i_λ , i_κ . Then we can multiply out the parentheses on both sides of the equation using the law $\overline{(\overline{XY})} = \overline{YX}$. It can accordingly be seen that, because of the non-commutativity of quaternion multiplication, equality can exist only if variables which are in juxtaposition on one side are also in juxtaposition on the other side. Since Z_θ and Z_i , as well as Z_λ and Z_κ , are in every case adjacent on the left side of the equation, they are again adjacent on the right side. Therefore $\overline{(Y^{(\tau\rho\sigma)} i_\nu^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)})}$ can not be equal to $\overline{(Y^{(\tau\rho\sigma)} i_\nu^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)})}$, and hence $\overline{(Y^{(\tau\rho\sigma)} i_\nu^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)})} \neq \overline{(Y^{(\tau\rho\sigma)} i_\nu^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)})}$. Similarly the invariable juxtaposition of Z_κ and Z_θ on the right side of the equation, as well as of Z_i and Z_λ ,

leads to the condition $(\nu\rho\sigma) \neq (\sigma\tau\nu)$. Now our equation assumes the form (since $R(XY) = R(YX)$)

$$R(Y^{(\tau\nu\sigma)} Y^{(\sigma\nu\rho)} i_{\theta}^{(\rho\sigma\nu)} Y^{(\nu\rho\sigma)} i_{\epsilon}^{(\sigma\nu\rho)} Y^{(\rho\sigma\nu)} Y^{(\nu\sigma\tau)} i_{\lambda}^{(\tau\nu\sigma)} Y^{(\sigma\tau\nu)} i_{\kappa}^{(\nu\sigma\tau)}) \\ = R(Y^{(\tau\nu\rho)} i_{\theta}^{(\rho\tau\nu)} Y^{(\nu\rho\tau)} Y^{(\tau\rho\sigma)} i_{\epsilon}^{(\sigma\tau\rho)} Y^{(\rho\sigma\tau)} i_{\lambda}^{(\tau\rho\sigma)} Y^{(\sigma\tau\rho)} Y^{(\rho\tau\nu)} i_{\kappa}^{(\nu\rho\tau)}).$$

Again we can replace i_{θ} , i_{ϵ} , i_{λ} , i_{κ} by Z_{θ} , Z_{ϵ} , Z_{λ} , Z_{κ} and it results that

$$(\rho\sigma\nu) = (\rho\tau\nu), (\sigma\nu\rho) = (\sigma\tau\rho), (\tau\nu\sigma) = (\tau\rho\sigma), (\nu\sigma\tau) = (\nu\rho\tau)$$

in addition to the conditions

$$(\tau\nu\rho) \neq (\rho\sigma\tau), (\nu\rho\sigma) \neq (\sigma\tau\nu)$$

already obtained, so that

$$(Y) \quad \begin{cases} Y^{(\tau\nu\sigma)} Y^{(\sigma\nu\rho)} = Y^{(\tau\nu\rho)}, & Y^{(\nu\rho\sigma)} = Y^{(\nu\rho\tau)} Y^{(\tau\rho\sigma)}, \\ Y^{(\rho\sigma\nu)} Y^{(\nu\sigma\tau)} = Y^{(\rho\sigma\tau)}, & Y^{(\sigma\tau\nu)} = Y^{(\sigma\tau\rho)} Y^{(\rho\tau\nu)}. \end{cases}$$

It is to be noted that we have required of the sequences ρ, σ, τ ; ρ, σ, ν ; ρ, τ, ν ; and σ, τ, ν that each of them be transformable into a monotonically increasing sequence by a cyclic permutation, and of course the same must hold also of ρ, σ, τ, ν .

We now see that the four $\bullet$ -equations simply comprise all possible subcases of $\alpha\epsilon\beta = \alpha\eta\beta$ for which α, ϵ, β and α, η, β are allowable sequences. The two next following $\bullet$ -inequalities express together that $\alpha\epsilon\beta \neq \beta\eta\alpha$, where α, ϵ, β and β, η, α are arbitrary allowable sequences. By the operation of (48 a) we can, for every pair $\alpha > \beta$, make a $\alpha\epsilon\beta$ vanish or make a $\beta\eta\alpha$ equal to $\overline{}$.²⁴ From our $\bullet$ -equations and $\bullet$ -inequalities it follows that in any case every $\alpha\epsilon\beta$ will vanish and every $\beta\eta\alpha$ will be equal to $\overline{}$. Hence, for $\rho < \sigma < \tau$, $\rho\sigma\tau$ equals $\overline{}$, while $\sigma\tau\rho$ and $\tau\rho\sigma$ vanish.

The Y -equations all arise by cyclic permutations of the first which may also be written

$$Y^{(\tau\nu\sigma)} = Y^{(\tau\nu\rho)} \overline{Y^{(\sigma\nu\rho)}}.$$

If now $\bar{\nu}$ is the immediate successor of ν (that is, $\nu + 1$ for $\nu \neq \gamma$, 1 for $\nu = \gamma$), then we have

$$Y^{(\tau\nu\sigma)} = T^{(\tau\nu)} \overline{T^{(\sigma\nu)}},$$

where we define $T^{(\omega\nu)} = Y^{(\omega\nu\bar{\nu})}$ ($\omega \neq \nu$) so long as the right side has meaning, and if it is meaningless ($\omega = \bar{\nu}$) then $T^{(\omega\nu)} = 1$. In this equation it is obvious that $\tau \neq \nu$; for $\tau \neq \bar{\nu}$ the equation follows from the preceding one with $\rho = \bar{\nu}$, and for $\tau = \bar{\nu}$ it is an identity. (Note that $T^{(\rho\sigma)}$ and $T^{(\sigma\rho)}$ are not identical.)

Among the operations of (48 b) we select

$$X_1^{(\alpha\beta)} = T^{(\beta\alpha)}, X_2^{(\alpha\beta)} = \overline{T^{(\alpha\beta)}} \quad (\alpha < \beta).$$

²⁴ We must leave both alternatives open since, for example, $\alpha\epsilon\beta$ is never an allowable sequence for $\alpha = \gamma$ and $\beta = 1$, and the same holds for β, η, α when $\alpha = \beta + 1$.

(The left sides are symmetric in α and β while the right sides are not. Hence, for $\alpha > \beta$, we obtain valid formulas by interchanging α and β on the right sides.) If $\rho < \sigma < \tau$, then, by (48 b), $Y^{(\rho\sigma\tau)}$, $Y^{(\sigma\tau\rho)}$, $Y^{(\tau\rho\sigma)}$ all go over into 1 since $\binom{\rho\sigma\tau}{\bullet\bullet\bullet} = \binom{\sigma\tau\rho}{\bullet\bullet\bullet} = \binom{\tau\rho\sigma}{\bullet\bullet\bullet} = 0$. We have therefore made all the Y 's simultaneously equal to 1, and if $\rho < \sigma < \tau$, then, by (47),

$$(49) \quad 2A_{\nu\lambda\mu}^{\rho\sigma\tau} = R(i_\nu i_\lambda i_\mu).$$

This result is exactly similar to that for $2A_{\nu\lambda\mu}^{123}$ in (45).

The multiplication rules (22) are accordingly uniquely determined by (49) and (23). Therefore in the preceding case there is at most one solution. But in any event a solution exists, namely, the Hermitian quaternion matrices of order γ (cf. paragraph 17). Hence this is the only solution.

20. $\gamma \geq 4$, $\chi = 2$. In (45) we have obtained a solution $2A_{\nu\lambda\mu}^{123}$ in the form $R(i_\nu i_\lambda i_\mu)$, where $i_1 = 1$ and $i_2 = i$ constitute the basis of the system of ordinary complex numbers. By (38) the most general solution arises if we replace

$$i_\nu \text{ by } \sum_{\nu'=1}^2 \psi_{\nu\nu'} i_{\nu'}, \quad i_\lambda \text{ by } \sum_{\lambda'=1}^2 \varphi_{\lambda\lambda'} i_{\lambda'}, \quad i_\mu \text{ by } \sum_{\mu'=1}^2 \omega_{\mu\mu'} i_{\mu'}.$$

(We make no use of the restrictions proved at the end of paragraph 15). But every orthogonal transformation $\sum_{\kappa'=1}^2 \sigma_{\kappa\kappa'} i_{\kappa'}$ of the basis of the complex number system may also be represented in the form $X i_\kappa \bullet$, where X is a complex number of absolute value unity and $\bullet$ represents nothing at all or else the sign for the complex conjugate. The most general expression for $2A_{\nu\lambda\mu}^{123}$ is therefore (note the commutativity of complex numbers)

$$2A_{\nu\lambda\mu}^{123} = R(Y^{(123)} i_\nu^{(312)} i_\lambda^{(123)} i_\mu^{(231)}),$$

where $Y^{(123)}$ is a certain complex number of absolute value unity and where the $\bullet$'s are independent of one another. It follows by (23) that this equation also holds for $2A_{\nu\lambda\mu}^{231}$ and $2A_{\nu\lambda\mu}^{312}$ if the indices 1, 2, and 3 are correspondingly permuted on the right side, provided that we define $Y^{(123)} = Y^{(231)} = Y^{(312)}$. Therefore we may set

$$(50) \quad 2A_{\nu\lambda\mu}^{\rho\sigma\tau} = R(Y^{(\rho\sigma\tau)} i_\nu^{(\tau\rho\sigma)} i_\lambda^{(\rho\sigma\tau)} i_\mu^{(\sigma\tau\rho)}), \quad Y^{(\rho\sigma\tau)} = Y^{(\sigma\tau\rho)} = Y^{(\tau\rho\sigma)},$$

where the Y 's are certain complex numbers of absolute value unity and where the $\bullet$'s are independent of one another. As in paragraph 19 we limit ourselves to such values of ρ, σ, τ which may be transformed into a monotonically increasing sequence by a cyclic permutation.

Again we must determine the influence of the basis transformations of the $M^{\rho\sigma}$ (in the sense of Theorem 13, (21) and (37)) on (50). This is exhibited in the formulas (47) for $2A_{\nu\lambda\mu}^{\rho\sigma\tau}$, $2A_{\nu\lambda\mu}^{\sigma\tau\rho}$, $2A_{\nu\lambda\mu}^{\tau\rho\sigma}$, or $2A_{\nu\lambda\mu}^{\rho\sigma\tau}$, $2A_{\nu\lambda\mu}^{\sigma\tau\rho}$, $2A_{\nu\lambda\mu}^{\tau\rho\sigma}$ as a transformation of the corresponding i_μ, i_λ, i_ν and hence as that of the i_κ with $\binom{\sigma\tau\rho}{\bullet\bullet\bullet}$ or $\binom{\rho\sigma\tau}{\bullet\bullet\bullet}$.

If we transform all the $M^{\rho\sigma}$ simultaneously so that i_κ in $M^{\rho\sigma}$ goes over into $X^{(\rho\sigma)} i_\kappa^{(\rho\sigma)}$, and if for the sake of simplicity we identify $\rho\sigma$ and $\sigma\rho$, then we have (note the commutativity of complex numbers)

$$(51) \quad \begin{cases} (a) & (\rho\sigma\tau) \text{ is replaced by } (\rho\sigma\tau)^{(\rho\sigma\tau)}, \\ (b) & Y^{(\rho\sigma\tau)} \text{ is replaced by } Y^{(\rho\sigma\tau)} X^{(\sigma\tau)} X^{(\tau\rho)} X^{(\rho\sigma)} X^{(\sigma\tau\rho)}. \end{cases}$$

Further, since $R(X) = R(\bar{X})$ and $\bar{X}\bar{Y} = \bar{X}\bar{Y}$, we can immediately

$$(51) \quad \begin{cases} (c) & \text{alter all } (\rho\sigma\tau), (\sigma\tau\rho), \text{ and } (\tau\rho\sigma) \text{ simultaneously,} \\ & \text{and replace } Y^{(\rho\sigma\tau)} \text{ by } \overline{Y^{(\rho\sigma\tau)}} \text{ in (50).} \end{cases}$$

By application of (51 c) we can simultaneously obtain $(\rho\sigma\tau) = \overline{}$ for all ρ, σ, τ such that $\rho < \sigma < \tau$.

We now consider (25). We assume $\rho < \sigma < \tau < \nu$. Then it follows that $(\rho\sigma\tau) = (\rho\sigma\nu) = (\rho\tau\nu) = (\sigma\tau\nu) = \overline{}$. If we substitute

$$\sum_{\kappa=1}^2 R(X i_\kappa^{(1)}) R(Y i_\kappa^{(2)}) = R(X \bar{Y}^{(1)}^{(2)})$$

and the particular choice of some $\bullet$'s just now formulated into (50), we get

$$\begin{aligned} & R(Y^{(\rho\sigma\nu)} Y^{(\sigma\tau\nu)} (v\rho\sigma) i_\theta i_i^{(u\rho\sigma)} i_\lambda^{(v\rho\sigma)} i_\kappa^{(v\rho\sigma)} i_\nu^{(v\rho\sigma)}) \\ &= R(Y^{(\rho\tau\nu)} Y^{(\rho\sigma\tau)} (v\rho\tau) i_\kappa^{(v\rho\tau)} i_\theta i_i^{(v\rho\tau)} i_\lambda^{(v\rho\tau)} i_\nu^{(v\rho\tau)}). \end{aligned}$$

Since we are forming linear aggregates we can set free complex variables $Z_\theta, Z_i, Z_\lambda, Z_\kappa$ in place of $i_\theta, i_i, i_\lambda, i_\kappa$. The above equality can hold only if the two parentheses are equal or if one is equal to the conjugate of the other. But since $\bar{Z}_\theta$ occurs on both sides the latter possibility is excluded. Consequently it must be the case that

$$(v\rho\sigma) = (\sigma\tau\rho)(v\rho\tau), (v\rho\tau)(v\rho\sigma) = (\tau\rho\sigma)(v\rho\tau), (v\rho\tau)(v\rho\sigma) = (v\rho\tau).$$

The first equation can also be written $(v\rho\sigma) = (\sigma\tau\rho)(\sigma\rho\nu)$ so that if $\bar{\rho}$ is the immediate successor of ρ (that is, $\rho + 1$ inasmuch as $\rho \neq \gamma$ because of the condition $\rho < \sigma < \tau < \nu$)

$$(v\rho\sigma) = ((\tau\rho))((v\rho)),$$

where we define $((\omega\rho)) = (\bar{\omega}\omega\rho)$ ($\omega \neq \rho$) so long as the right side has meaning, and if it is meaningless ($\omega = \bar{\rho}$) then $((\omega\rho)) = 0$.

For $\rho < \sigma < \tau$ it follows that $(\rho\sigma\tau) = \overline{}$ and $(\sigma\tau\rho) = ((\sigma\rho))((\tau\rho))$, where $((\tau\rho))$ is arbitrary. If for $((\tau\rho)) = \overline{}$ we apply the operation (51 c) and for $((\tau\rho)) = 0$ we do nothing, then it results that $(\rho\sigma\tau) = \overline{((\tau\rho))}$ and $(\sigma\tau\rho) = ((\sigma\rho))$, where $(\tau\rho\sigma)$ is arbitrary. If to this result we apply the operation (51 a) with $(\alpha\beta) = ((\alpha\beta))$ ($\alpha > \beta$), then it results that $(\rho\sigma\tau) = \overline{}$ and $(\sigma\tau\rho) = 0$, where $(\tau\rho\sigma)$ is arbitrary. From these equations we have

$$(v\rho\sigma) = (\tau\rho\sigma), (v\rho\tau)(v\rho\sigma) = (v\rho\tau).$$

The first states that $(\omega\rho\sigma)$ (naturally $\rho < \sigma < \omega$) is independent of ω so that $(\omega\rho\sigma) = ((\rho\sigma))$ and the second states that $((\sigma\tau)) ((\rho\sigma)) = ((\rho\tau))$. Therefore $(((\rho\sigma))) = [\rho][\sigma]$ and $(\omega\rho\sigma) = [\rho][\sigma]$. Therefore $(\rho\sigma\tau) = \overline{\quad}$, $(\sigma\tau\rho) = 0$, and $(\tau\rho\sigma) = [\rho][\sigma]$. If for $[\rho] = \overline{\quad}$ we here apply the operation (51 c) and for $[\rho] = 0$ we do nothing, then it results that $(\rho\sigma\tau) = [\rho]$, $(\sigma\tau\rho) = [\rho]$, and $(\tau\rho\sigma) = [\sigma]$. If we apply the operation (51 a) with $(\alpha\beta) = [\beta]$ ($\alpha > \beta$), then we have $(\rho\sigma\tau) = \overline{\quad}$ and $(\sigma\tau\rho) = (\tau\rho\sigma) = 0$.

In the second equation following (51 c) the coefficients in the parentheses must be equal to one another so that for $\rho < \sigma < \tau < \nu$ it follows that $Y^{(\rho\sigma\nu)} Y^{(\sigma\tau\nu)} (\nu\rho\sigma) = Y^{(\rho\tau\nu)} Y^{(\rho\sigma\tau)} (\tau\nu\rho)$ so that with the $\cdot$ just obtained, $Y^{(\rho\sigma\nu)} Y^{(\sigma\tau\nu)} = Y^{(\rho\tau\nu)} Y^{(\rho\sigma\tau)}$ and hence $Y^{(\rho\sigma\tau)} = Y^{(\sigma\tau\nu)} \overline{Y^{(\rho\tau\nu)}} Y^{(\rho\sigma\nu)}$. Consequently

$$Y^{(\rho\sigma\tau)} = T^{(\sigma\tau)} \overline{T^{(\rho\tau)}} T^{(\rho\sigma)},$$

where we set $T^{(\alpha\beta)} = Y^{(\alpha\beta\gamma)}$ ($\alpha < \beta$) so long as the right side has meaning, and if it is meaningless ($\beta = \gamma$) then $T^{(\alpha\beta)} = 1$. We now select from the operations of (51 b)

$$X^{(\alpha\beta)} = \overline{T^{(\alpha\beta)}} \quad (\alpha < \beta).$$

(The left side is symmetric in α and β while the right side is not so that we obtain the formulas valid for $\alpha > \beta$ by interchange of α and β on the right side.) Then all $Y^{(\rho\sigma\tau)}$ go over into 1.

By (50) it follows that, for $\rho < \sigma < \tau$,

$$(52) \quad 2A_{\nu\lambda\mu}^{\rho\sigma\tau} = R(i, \overline{i_\lambda} i_\mu)$$

which is similar to (49) and to $2A_{\nu\lambda\mu}^{123}$ in (45).

Hence the multiplication rules (22) are uniquely determined by (52) and (23). There is therefore at most one solution in the preceding case. But a solution exists in any event, for example, the Hermitian complex matrices of order γ (cf. paragraph 17). Hence it is the only solution.

21. $\gamma \geq 4$, $\chi = 1$. In this case it is not worth while to carry out the general theory of paragraphs 14 to 17. Since $\lambda = \mu = \nu = 1$ in $A_{\nu\lambda\mu}^{\rho\sigma\tau}$, therefore $2A_{\nu\lambda\mu}^{\rho\sigma\tau} = Y^{\rho\sigma\tau}$ and (24) leads to the result $Y^{\rho\sigma\tau} = \pm 1$. The basis transformations of $M^{\rho\sigma}$ in the sense of Theorem 13, (21), and (37) are $(s_1^{\rho\sigma})' = X^{\rho\sigma} s_1^{\rho\sigma}$ and $X^{\rho\sigma} = \pm 1$, where for the sake of simplicity we identify $\rho\sigma$ and $\sigma\rho$. Hence

$$(53) \quad Y^{(\rho\sigma\tau)} \text{ is replaced by } Y^{(\rho\sigma\tau)} X^{(\sigma\tau)} X^{(\tau\rho)} X^{(\rho\sigma)}.$$

In the preceding case it follows from (25) that

$$(54) \quad Y^{(\rho\sigma\nu)} Y^{(\sigma\tau\nu)} = Y^{(\rho\tau\nu)} Y^{(\rho\sigma\tau)}.$$

The relation

$$Y^{(\rho\sigma\tau)} = T^{(\sigma\tau)} T^{(\rho\tau)} T^{(\rho\sigma)}$$

follows from (54) in exactly the same way as did the analogous relation for the complex Y at the end of paragraph 20. We similarly find that every T is either a Y or 1 and hence is ± 1 . Therefore we select from the operations of (53)

$$X^{(\alpha\beta)} = T^{(\alpha\beta)} \quad (\alpha < \beta).$$

(With reference to the symmetry conditions in α and β , cf. the remark at the end of paragraph 20.) Then all $Y^{(\rho\sigma\tau)}$ go over into 1. Hence we have in general that

$$(54) \quad 2A_{\nu\mu}^{\rho\sigma\tau} = 1.$$

The multiplication rules (22) are uniquely determined by (54) and (23). Hence in the preceding case there is at most one solution. But in any event a solution exists, for example, the Hermitian (real) matrices of order γ (cf. paragraph 17). Hence it is the only solution.

22. We now summarize our results.

In the course of the preceding investigations we have learned the properties of the following algebras:

α . $\mathfrak{R}$, the algebra of the real numbers in which $a \pm b$, λa , and ab are defined in the usual way.

β . $\mathfrak{S}_N$, $N = 1, 2, \dots$, the algebra with the linear basis $1, s_1, \dots, s_{N-1}$, in which $a \pm b$ and λa are defined in the usual way but in which ab is defined by

$$11 = 1, 1s_\mu = s_\mu, \quad s_\mu s_\nu = \delta_{\mu\nu} 1.$$

γ . $\mathfrak{M}_\gamma^\chi$, $\chi = 1, 2, 4, 8$; $\gamma = 1, 2, \dots$, the algebra of the Hermitian matrices of order γ in which the elements are real numbers, ordinary complex numbers, quaternions, and quasi-quaternions respectively for $\chi = 1, 2, 4$, and 8. $a \pm b$ and λa are defined in the usual way, but ab is defined by $\frac{a \times b + b \times a}{2}$, where $a \times b$ represents the usual matrix multiplication.

We can summarize the results of paragraphs 12 to 21 as follows:

FUNDAMENTAL THEOREM 2. *The only irreducible r -number algebras are the following:*

For $\gamma = 1$: $\mathfrak{R}$.

For $\gamma = 2$: All $\mathfrak{S}_N$, $N = 3, 4, \dots$. ($\mathfrak{S}_1$ is $\mathfrak{R}$ and $\mathfrak{S}_2$ is reducible.)

For $\gamma = 3$: All $\mathfrak{M}_3^\chi$, $\chi = 1, 2, 4, 8$.

For $\gamma \geq 4$: All $\mathfrak{M}_\gamma^\chi$, $\chi = 1, 2, 4$. (All $\mathfrak{M}_1^\chi$ are $\mathfrak{R}$; the $\mathfrak{M}_2^\chi$ are respectively the $\mathfrak{S}_{\chi+2}$; the $\mathfrak{M}_3^\chi$ are already classified under $\gamma = 3$; the $\mathfrak{M}_\gamma^8$, $\gamma \geq 4$, are not r -number algebras.)

In connection with this result we mention the following circumstance: let A' be an associative, but not necessarily commutative algebra having the real

numbers as coefficients. Then $a \pm b$, λa , and $a \times b$ are the fundamental operations and all the usual rules of calculation are valid except for $a \times b = b \times a$.

We now define $ab = \frac{a \times b + b \times a}{2}$ and we consider a subsystem A of A' which is closed for $a \pm b$, λa , and ab (but not necessarily for $a \times b$) and which is "formally real" ($a^2 = a \times a = aa$; cf. (I) in paragraph 1 and footnote 5). Since (II) is obviously satisfied (a^n is the same for ab and $a \times b$), A is an r -number algebra. By the theorem of Wedderburn²⁵ we could select a matrix algebra for A' without loss of generality. We say (in agreement with the terminology of footnote 1) that an r -number algebra which is isomorphic with such an A results from quasi-multiplication. As one of us has shown elsewhere (cf. footnote 1), essentially new results, in contrast with the present content of quantum mechanics, are to be expected only in such r -number algebras which do not result from quasi-multiplication.

Now all the r -number algebras of our list, with the exception of $\mathfrak{M}_3^8$, result from quasi-multiplication. This is obvious in the case of $\mathfrak{R}$ and the $\mathfrak{M}_\chi^x$ ($\chi = 1, 2, 4$) so that only the $\mathfrak{S}_N$ need examination. But in any event the $\mathfrak{S}_N$ arise by the operation $ab = \frac{a \times b + b \times a}{2}$ from a system in which 1 is the modulus, $s_\mu^2 = 1$, and $s_\mu s_\nu = -s_\nu s_\mu$ ($\mu \neq \nu$, $\mu, \nu = 1, \dots, N - 1$). We now take the quaternion-like, non-commutative, but associative system 1, i, j, k with $i^2 = j^2 = k^2 = 1$, $ij = -ji = k$, $jk = -kj = i$, $ki = -ik = j$ and form $N - 1$ sets $1^{(1)}, i^{(1)}, j^{(1)}, k^{(1)}; \dots; 1^{(N-1)}, i^{(N-1)}, j^{(N-1)}, k^{(N-1)}$ of this sort. The direct products $1 = 1^{(1)} \cdot 1^{(2)} \cdot \dots \cdot 1^{(N-1)}$, $s_1 = j^{(1)} \cdot 1^{(2)} \cdot \dots \cdot 1^{(N-1)}$, $s_2 = i^{(1)} \cdot j^{(2)} \cdot \dots \cdot 1^{(N-1)}$, $\dots$, $s_{N-1} = i^{(1)} \cdot i^{(2)} \cdot \dots \cdot j^{(N-1)}$ perform the desired function.

On the other hand it can be shown that $\mathfrak{M}_3^8$ can not result from quasi-multiplication. The proof of this remarkable result was given by A. Albert, to whom we proposed the problem, and can be found in his next paper in this journal. Hence there is one single (irreducible) r -number algebra which does not arise from quasi-multiplication: $\mathfrak{M}_3^8$.

²⁵ Proc. London Math. Soc., 1907, Vol. (2) 6, p. 99.

Zum Haarschen Maß in topologischen Gruppen

1. G sei eine topologische Gruppe, d.h. eine, in der ein Hausdorffscher Umgebungsbegriff definiert ist ¹⁾, so daß die fundamentalen Gruppen-Operationen $a \cdot b$ und a^{-1} stetige Funktionen von a, b bzw. a sind. Ferner werde G stets als separabel vorausgesetzt ²⁾. G heißt bekanntlich kompakt, wenn jede unendliche Teilmenge von G mindestens einen Häufungspunkt besitzt ³⁾, und lokal kompakt, wenn jeder Punkt von G eine Umgebung hat, deren abgeschlossene Hülle kompakt ist (es genügt natürlich, dies für die Einheit 1 zu postulieren: die topologische Abbildung $x \rightarrow x \cdot a$ führt sie ja in ein beliebiges a über). In der gruppentheoretischen Terminologie heißen kompakte G auch geschlossen, und nicht-kompakte offen.

A. HAAR bewies ⁴⁾, daß in jedem im kleinen kompakten G ein Maßbegriff definiert werden kann, der alle formalen Eigenschaften des Lebesgueschen Maßes besitzt ⁵⁾, und gegenüber jeder Abbildung $x \rightarrow x \cdot a$ invariant ist. Dabei kann das Maß von ganz G durchaus unendlich sein, aber es gilt: Jedes Maß ist ≥ 0 , jedes Maß einer offenen (nicht-leeren) Menge ist > 0 , jedes Maß einer kompakten Menge ist endlich.

¹⁾ Für die Grundbegriffe der Topologie vgl. etwa HAUSDORFF, Mengenlehre, Berlin u. Leipzig (1927), § 40, S. 226—230.

²⁾ Wir verlangen somit von HAUSDORFFS a.a.O. aufgezählten Axiomen 1, 2, 3, 6, 10, aus diesen folgt das u.U. ebenfalls erwünschte 8. Vgl. a.a.O. S. 230 oben.

³⁾ Vgl. a.a.O. S. 107.

⁴⁾ Annals of Math. 34 (1933), 147—169 (§ 3).

⁵⁾ Dieselben werden von CARATHEODORY in seinem Buche „Reelle Funktionen“, Berlin u. Leipzig (1918), 237—243 abstrakt diskutiert; für beliebige topologische Räume vgl. auch die Arbeit des Verf. [Annals of Math. 33 (1932), 572—586, Def. 2, 4 auf S. 574, 576].

Ein solches Maß nennen wir ein Haarsches rechts-invariantes Maß. An Stelle der Invarianz gegenüber $x \rightarrow x \cdot a$ kann auch jene gegenüber $x \rightarrow a \cdot x$ erreicht werden, dann nennen wir es ein Haarsches links-invariantes Maß.

Es ist zu vermuten, daß es in einem gegebenen G , bis auf das triviale Multiplizieren aller Maße mit einem gemeinsamen (konstanten) positiven Faktor, nur *ein* Haarsches rechts- (bzw. links-) invariantes Maß gibt. Diese Frage konnte jedoch bisher nicht allgemein entschieden werden. Im Folgenden soll u.a. gezeigt werden, daß sie für kompakte G zu bejahen ist.

Wir werden nämlich für kompakte G eine neue Methode angeben, ein Haarsches rechts-invariantes Maß aufzustellen — das übrigens in diesem Falle von selbst auch links-invariant sein wird ⁶⁾, ja auch gegenüber der Abbildung $x \rightarrow x^{-1}$ ⁷⁾. Unsere Methode ist von der Haarschen wesentlich verschieden und vielleicht auch an und für sich nicht uninteressant; sie ergibt für kompakte G das Endresultat rascher, und so daß die oben erwähnte Eindeutigkeit von selbst mit herauskommt. Schließlich sei noch erwähnt, daß dieses Resultat, obwohl spezieller als das Haarsche, zur Begründung des von Verf. bewiesenen Satzes, wonach jede geschlossene endlich-viel-parametrische Gruppe eine Liesche Gruppe ist, sowie der weiteren hieran anschließenden Sätze ⁸⁾, ausreicht. Genauer: es wird für diese Anwendung nur das in 2 definierte und in 3 konstruierte „Mittel stetiger Funktionen“ gebraucht, und nicht das Haar-Lebesguesche Maß selbst — d.h. die in 2 gegebene Herleitung des Letzteren aus dem Ersteren ist entbehrlich.

2. An Stelle des Haar-Lebesgueschen Maßes werden wir ein Mittel stetiger Funktionen definieren, d.h. Folgendes:

Jeder reellwertigen stetigen, in G definierten Funktion $f(x)$ wird eine reelle Zahl $M(f(x))$ zugeordnet, so daß

1. $M(\alpha f(x)) = \alpha M(f(x))$ (α eine reelle Zahl).
2. $M((f(x) + g(x))) = M(f(x)) + M(g(x))$.
3. Wenn stets $f(x) \geq 0$ ist, so ist $M(f(x)) \geq 0$.
4. Wenn stets $f(x) = 1$ ist, so ist $M(f(x)) = 1$.

⁶⁾ Für nicht-kompakte G ist dies nicht immer der Fall.

⁷⁾ Für Liesche Gruppen gaben F. PETER und H. WEYL ein solches Maß an [Math. Ann. 97 (1928), 737—755 (737)]. Aber wir gehen, Haar folgend, rein mengentheoretisch, ohne jede Regularitäts-Annahme vor.

⁸⁾ A. *Annals of Math.* 34 (1933), 170—179 (Satz 1, S. 182, und Satz 2, S. 187).

5. $M(f(x \cdot a)) = M(f(x))$
 6. $M(f(a \cdot x)) = M(f(x))$
 7. $M(f(x^{-1})) = M(f(x))$.
- $$\left. \begin{array}{l} 5. \\ 6. \end{array} \right\} \quad (a \text{ ein Element von } G).$$

Mit der Hilfe eines solchen Mittels kann nämlich ein Haar-Lebesguesches Maß eingeführt werden, wie die folgenden Überlegungen zeigen:

A. Sei O eine offene Teilmenge von G . Wir schreiben $f(x) \subset O$ ($f(x)$ stetig!), falls $f(x)$ in $O \geq 0, \leq 1$ und außerhalb von $O = 0$ ist. Die obere Grenze aller $M(f(x))$, $f(x) \subset O$ (diese sind ja nach 3, 4 alle $\geq 0, \leq 1$) heiße $M(O)$.

B. Sei O Teilmenge von $\sum_{n=1}^{\infty} O_n$ (der Fall endlich vieler Addenden kann mit eingeschlossen werden, wenn man $O_{m+1}, O_{m+2}, \dots$ leer ansetzt) dann ist

$$M(O) \leq \sum_{n=1}^{\infty} M(O_n).$$

Sei nämlich $f(x) \subset O$. Die Menge A_ε , $\varepsilon > 0$, der x mit $f(x) \geq \varepsilon$ ist abgeschlossen. Konstruieren wir für jedes n eine Folge offener Mengen $O_n^1 \subset O_n^2 \subset \dots$, die mitsamt ihren abgeschlossenen Hüllen $\bar{O}_n^1 \subset \bar{O}_n^2 \subset \dots \subset O_n$ sind und O_n zur Vereinigungsmenge haben⁹). Wegen der Kompaktheit folgt aus $A_\varepsilon \subset O \subset \sum_{n=1}^{\infty} O_n = \sum_{n=1}^{\infty} \bar{O}_n^m$, daß A_ε schon in der Summe endlich vieler $\bar{O}_n^m$ enthalten ist. Indem wir für jedes n nur jenes mit größtem m beibehalten, erreichen wir, daß kein n mehr als einmal vorkommt — und, wenn das größte $n = N$ ist, können wir, durch Hinzufügen eines beliebigen $\bar{O}_n^m$ für jedes fehlende $n \leq N$, erreichen, daß genau die $n \leq N$ vorkommen. Also: $A_\varepsilon \subset \bar{O}_1^m + \dots + \bar{O}_N^m$. Da $\bar{O}_n^m \subset O_n$ ist, existiert eine stetige Funktion $f_n(x)$, die stets $\geq 0, \leq 1$, in $\bar{O}_n^m = 1$ und außerhalb $O_n = 0$ ist¹⁰). In A_ε haben wir also $f(x) \leq 1$, mindestens ein $f_n(x) = 1$, $n = 1, \dots, N$,

⁹) Daß das geht, folgt bekanntlich aus dem Axiom 8, a.a.O. Anm. 2). Liegt ein Entfernungsbegriff in G vor, so ist es evident.

¹⁰) Dieser Satz stammt von P. URYSOHN [Math. Ann. 94 (1925), 309]; vgl. auch K. Menger, Dimensionstheorie [Berlin u. Leipzig (1928), 57—59].

also $\sum_1^N f_n(x) \geq 1$; außerhalb A_ε $f(x) < \varepsilon$, $\sum_1^N f_n(x) \geq 0$. Also ist jedenfalls $f(x) < \sum_1^N f_n(x) + \varepsilon$. Ferner ist $f_n(x) \subset O_n$.

Aus alledem folgt nach 1—4

$$M(f(x)) \leq \sum_1^N M(f_n(x)) + \varepsilon \leq \sum_1^N M(O_n) + \varepsilon \leq \sum_1^\infty M(O_n) + \varepsilon.$$

Da dies für alle $f(x) \subset O$ gilt, ist auch $M(O) \leq$ als die rechte Seite, und da es für jedes $\varepsilon > 0$ gilt, können wir ε fortlassen, d.h. es ergibt sich die Behauptung.

C. Jedes $M(O)$ ist ≥ 0 , ≤ 1 ; für $O = G$ ist es $= 1$; ist O nicht leer, so ist es > 0 . Die zwei ersten Behauptungen sind klar. Wäre die dritte falsch, so wäre $M(O) = 0$, und wenn O_a das $(x \rightarrow x \cdot a)$ -Bild von O ist, $M(O_a) = 0$. Gehört b_0 zu O , so gehört b zu $O_{b_0^{-1}b}$, also überdecken die offenen O_a zusammen ganz G . Also überdecken es wegen der Kompaktheit schon endlich viele ¹¹⁾, und da deren $M(O_a) = 0$ sind, wäre $M(G) \leq 0$, d.h. $= 0$, was unmöglich ist.

D. $M(O)$ ist gegenüber einer jeden der Abbildungen $x \rightarrow x \cdot a$, $x \rightarrow a \cdot x$, $x \rightarrow x^{-1}$ invariant. Dies folgt unmittelbar aus 5—7.

Auf Grund von A—D können wir nun das Haar-Lebesguesche Maß auf die beim Lebesgueschen Maße übliche Weise definieren: ist M eine beliebige Teilmenge von G , so ist $\mu^*(M)$ die untere

Grenze aller $\sum_1^\infty M(O_n)$, $M \subset \sum_1^\infty O_n$. Aus A—C folgt, daß die üblichen Schlußweisen beim Lebesgueschen Maße (vgl. a.a.O. Anm. 5)) wörtlich auf $\mu^*(M)$ übertragen werden können, und aus D, daß es gegenüber $x \rightarrow x \cdot a$, $x \rightarrow a \cdot x$, $x \rightarrow x^{-1}$ invariant ist.

3. Wir konstruieren nunmehr das Mittel stetiger Funktionen im Sinne von 2.

Für jedes stetige $f(x)$ definieren wir die Schwankung $S(f(x))$ durch

$$\text{Max}(f(x)) - \text{Min}(f(x)) = \text{Max}(|f(x) - f(y)|) = S(f(x)).$$

Wir betrachten nun alle Funktionen $\frac{\sum_1^N f(x \cdot a_n)}{N}$, wobei N jede der Zahlen 1, 2, ... sein kann und $a_1, \dots, a_N$ beliebige Elemente von G sind.

¹¹⁾ Es würde genügen, aus der Separabilität auf abzählbar viele zu schließen.

$f(x)$ ist gleichmäßig stetig: d.h. es existiert zu jedem $\varepsilon > 0$ eine Umgebung V_ε der Einheit 1, so daß immer $|f(x) - f(y)| < \varepsilon$ ist, wenn $x \cdot y^{-1}$ zu V_ε gehört. Denn andernfalls könnten wir für jedes V^n einer sich auf 1 zusammenziehenden Umgebungsfolge $V^1, V^2, \dots$ zwei x_n, y_n^{-1} mit $|f(x_n) - f(y_n)| \geq \varepsilon, x_n y_n^{-1}$ in V^n angeben. Dann würde $x_n y_n^{-1}$ gegen 1 konvergieren. Die Folge $x_1, x_2, \dots$ hat, da G kompakt ist, einen Häufungspunkt $\bar{x}$, gegen den eine geeignete Teilfolge $x_{n_\nu}, \nu = 1, 2, \dots$ konvergiert. Da $x_{n_\nu} y_{n_\nu}^{-1}$ gegen 1 konvergiert, konvergiert auch y_{n_ν} gegen $\bar{x}$. Also konvergieren $f(x_{n_\nu}), f(y_{n_\nu})$ gegen $f(\bar{x}), |f(x_{n_\nu}) - f(y_{n_\nu})|$ gegen 0, im Widerspruch dazu, daß es für alle $\nu \geq \varepsilon$ sein sollte.

Alle $f(x \cdot a)$ sind mit $f(x)$ gleichartig stetig, d.h. für $x \cdot y^{-1}$ in V_ε gilt $|f(x \cdot a) - f(y \cdot a)| < \varepsilon$ mit demselben V_ε wegen $(x \cdot a) \cdot (y \cdot a)^{-1} = xy^{-1}$. Und daher gilt dies auch für alle

$$\frac{\sum_1^N f(x \cdot a_n)}{N}$$
 . Hieraus folgt auf Grund bekannter Schlußweisen,

daß diese ganze Funktionenklasse im Sinne der gleichmäßigen Konvergenz kompakt ist, d.h. daß aus jeder Folge ihrer Funktionen eine gleichmäßig konvergente Teilfolge ausgesondert werden kann¹²⁾.

Sei nun $\bar{s}$ die untere Grenze aller Schwankungen $S\left(\frac{\sum_1^N f(x \cdot a_n)}{N}\right)$ (für beliebige $N, a_1, \dots, a_N$), da diese alle ≥ 0 sind, ist $\bar{s} \geq 0$.

Es gibt also eine Folge von Funktionen von der Form $\frac{\sum_1^N f(x \cdot a_n)}{N}$,

deren Schwankungen gegen $\bar{s}$ konvergieren, und nach dem vorher Gesagten ist eine geeignete Teilfolge dieser Funktionenfolge gleichmäßig konvergent. Die Limesfunktion heiße $g(x)$, sie ist also stetig und hat die Schwankung $\bar{s}$. Ferner folgt aus ihrem Limes-

Charakter, daß sie durch Funktionen von der Form $\frac{\sum_1^N f(x \cdot a_n)}{N}$

beliebig gut (gleichmäßig) approximiert werden kann.

Nehmen wir an, ein $\frac{\sum_1^M g(x \cdot b_m)}{M}$ hätte eine kleinere Schwan-

¹²⁾ Vgl. hierzu z.B. M. FRÉCHET [Rend. Circolo Palermo 22 (1906), 1—72 („Théorèmes“, p. 13, 30)].

kung als $g(x)$. Etwa um $\varepsilon > 0$. Dann approximieren wir $g(x)$

mit einem $\frac{\sum_{n=1}^N f(x \cdot a_n)}{N}$ um $\leq \frac{\varepsilon}{3}$, wodurch auch $\frac{\sum_{m=1}^M g(x \cdot b_m)}{M}$ von

$\frac{\sum_{m=1}^M \sum_{n=1}^N f(x \cdot b_m \cdot a_n)}{MN}$ um $\leq \frac{\varepsilon}{3}$ approximiert wird. Daher unter-

scheiden sich ihre Schwankungen um $\leq 2\frac{\varepsilon}{3}$, d.h. die Schwankung der letzteren Funktion ist $\leq \bar{s} - \varepsilon + 2\frac{\varepsilon}{3} = \bar{s} - \frac{\varepsilon}{3} < s$. Dies ist unmöglich, da jene Schwankung $\geq \bar{s}$ sein muß.

Da $\frac{\sum_{m=1}^M g(x \cdot b_m)}{M}$ ein Max $\leq$ und ein Min $\geq$ hat als $g(x)$, ist

seine Schwankung $\leq$, da aber, wie gezeigt wurde, $<$ unmöglich ist, ist die Schwankung, und also auch Max und Min, dieselbe wie bei $g(x)$. Aus der Gleichheit der Max folgt aber: wenn

$\frac{\sum_{m=1}^M g(x \cdot b_m)}{M}$ sein Max an der Stelle $\bar{x}$ annimmt, so muß $g(x)$ sein

Max an allen Stellen $\bar{x} \cdot b_1, \dots, \bar{x} \cdot b_M$ gleichzeitig annehmen.

Sei nun $\bar{b}_1, \bar{b}_2, \dots$ eine in G überall dichte Folge. Wir bilden für $\bar{b}_1, \dots, \bar{b}_M, M = 1, 2, \dots$ das obige $\bar{x}$: $\bar{x}_M$. Es ist $g(\bar{x}_M \cdot \bar{b}_m) = \text{Max } (g(x))$ für $m \leq M$. Die $\bar{x}_1, \bar{x}_2, \dots$ haben einen Häufungspunkt $\bar{x}$, also auch eine Teilfolge $\bar{x}_{M_\nu}, \nu = 1, 2, \dots$, die gegen $\bar{x}$ konvergiert. Ist a ein beliebiges Element von G , so existiert eine Teilfolge von $\bar{b}_1, \bar{b}_2, \dots$, etwa $\bar{b}_{P_\nu}, \nu = 1, 2, \dots$, die gegen $x^{-1} \cdot a$ konvergiert. Dabei kann man (z.B. durch Wiederholung der P_ν) erreichen, daß stets $P_\nu \leq M_\nu$ gilt. So wird $g(\bar{x}_{M_\nu} \cdot \bar{b}_{P_\nu}) = \text{Max } (g(x))$ und $\bar{x}_{M_\nu} \cdot \bar{b}_{P_\nu}$ konvergiert gegen $x \cdot (x^{-1} \cdot a) = a$. Also ist $g(a) = \text{Max } (g(x))$ für jedes a , d.h. $g(x)$ ist konstant.

Wir nennen eine Zahl α ein rechts-Mittel von $f(x)$, falls zu jedem $\varepsilon > 0$ ein $N = 1, 2, \dots$ und N Elemente $a_1, \dots, a_N$ von G existieren, so daß für alle x

$$* \quad \alpha - \varepsilon \leq \frac{1}{N} \sum_{n=1}^N f(x \cdot a_n) \leq \alpha + \varepsilon$$

gilt. In dieser Ausdrucksweise lautet unser soeben gewonnenes

Resultat: es gibt (mindestens) ein rechts-Mittel α von $f(x)^{13}$.

4. Analog können wir die links-Mittel β durch die Existenz von $M = 1, 2, \dots$ und $b_1, \dots, b_M$ zu jedem $\varepsilon > 0$ mit

$$** \quad \beta - \varepsilon \leq \frac{\sum_{m=1}^M f(b_m \cdot x)}{M} \leq \beta + \varepsilon$$

definieren. Da G eine topologische Gruppe bleibt, wenn wir ab, a^{-1} durch ba, a^{-1} ersetzen, hat es dieselben Eigenschaften wie das rechts-Mittel: es gibt (mindestens) ein links-Mittel von $f(x)$.

Wen α ein rechts- und β ein links-Mittel von $f(x)$ ist, so gelten * und **. Ersetzen wir in * x durch $b_m \cdot x$ und bilden $\frac{1}{M} \sum_1^M$ und in ** x durch $x \cdot a_n$ und bilden $\frac{1}{N} \sum_1^N$, so entsteht beidemale

$$\frac{\sum_{m=1}^M \sum_{n=1}^N f(b_m \cdot x \cdot a_n)}{MN}, \text{ das also } \geq \alpha - \varepsilon, \beta - \varepsilon \text{ und } \leq \alpha + \varepsilon,$$

$\beta + \varepsilon$ ist. Daraus folgt $\alpha - \varepsilon \leq \beta + \varepsilon, \beta - \varepsilon \leq \alpha + \varepsilon$, d.h. $|\alpha - \beta| \leq 2\varepsilon$. Dies gilt für alle $\varepsilon > 0$, so daß $\alpha = \beta$ sein muß. Wir können dieses Resultat auch so formulieren:

$f(x)$ hat genau ein rechts-Mittel und genau ein links-Mittel, und die beiden sind einander gleich. Ihr gemeinsamer Wert heie $M(f(x))$.

Nunmehr wollen wir zeigen, da dieses $M(f(x))$ ein Mittel stetiger Funktionen ist im Sinne der am Anfang von 2 gegebenen

¹³⁾ Hieraus folgt bereits die Eindeutigkeit des Haarschen Maes fr G , wenn man dasselbe als bekannt voraussetzt. Sei nmlich $\int f(x)dx$ das auf dasselbe begrndete Lebesguesche Integral, dann ist fr jedes rechts-Mittel α

$$\int (\alpha - \varepsilon) dx \leq \int \frac{\sum_{n=1}^N f(x \cdot a_n)}{N} dx = \frac{1}{N} \sum_{n=1}^N \int f(x \cdot a_n) dx = \frac{1}{N} \sum_{n=1}^N \int f(x) dx = \int f(x) dx \leq \int (\alpha + \varepsilon) dx$$

$$\alpha - \varepsilon \leq \frac{\int f(x) dx}{\int 1 dx} \leq \alpha + \varepsilon \quad \text{d.h.} \quad \alpha = \frac{\int f(x) dx}{\int 1 dx}.$$

Somit gibt es nur ein einziges rechts-Mittel α (was wir in 4 auch direkt einsehen werden). Aber dieses wurde in 2 ohne Bezugnahme auf das Haarsche Integral $\int f(x)dx$ definiert, also legt es dieses (bis auf den konstanten Faktor $\int 1 dx$) eindeutig fest. Da also alle Haarschen Integrale fr stetige $f(x)$ dasselbe $\int f(x)dx$ ergeben, stimmen auch die Mabegriffe berein (alles bis auf den konstanten Faktor!)

Definition, d.h. daß es die dortigen Bedingungen 1 — 7 erfüllt. Daß es kein anderes Mittel stetiger Funktionen geben kann, ist klar: denn wäre $M'(f(x))$ eines, so ergäbe die Methode von Anm. 13) sofort $M'(f(x)) = M(f(x))$.

1, 3, 4 sind offenbar erfüllt, 5 ist evident, wenn man $M(f(x))$ als links-Mittel, 6 wenn man es als rechts-Mittel ansieht (man ersetze x durch $x \cdot a$ bzw. $a \cdot x$). Betrachten wir nun 2. Wir setzen $M(f(x)) = \alpha$, $M(g(x)) = \beta$. Wenn wir in $*$ x durch $b'_m \cdot x$

ersetzen und $\frac{1}{M'} \sum_1^{M'} f(b'_m \cdot x)$ bilden, so zeigt sich, daß auch $\frac{\sum_1^{M'} f(b'_m \cdot x)}{M'}$ das rechts-Mittel α hat, also auch das links-Mittel α . Da $g(x)$ das links-Mittel β hat, können wir

$$\beta - \frac{\varepsilon}{2} \leq \frac{\sum_1^M g(b_m \cdot x)}{M} \leq \beta + \frac{\varepsilon}{2}$$

erreichen. Nun sei $M' = M$, $b'_m = b_m$, da $\frac{\sum_1^M f(b_m \cdot x)}{M}$ das links-Mittel α hat, können wir auch

$$\alpha - \frac{\varepsilon}{2} \leq \frac{\sum_1^M \sum_1^N f(b_m \cdot a_n \cdot x)}{MN} \leq \alpha + \frac{\varepsilon}{2}$$

erreichen. Nun ersetzen wir in der ersten Gleichung x durch $a_n \cdot x$, bilden $\frac{1}{N} \sum_1^N$ und addieren hierzu die zweite Gleichung.

So wird

$$\alpha + \beta - \varepsilon \leq \frac{\sum_1^M \sum_1^N (f(b_m \cdot a_n \cdot x) + g(b_m \cdot a_n \cdot x))}{MN} \leq \alpha + \beta + \varepsilon$$

Somit ist $\alpha + \beta$ das links-Mittel von $f(x) + g(x)$ und 2 ist bewiesen.

Somit erfüllt $M(f(x))$ 1–6, woraus unmittelbar folgt, daß $M'(f(x)) = M(f(x^{-1}))$ 1–6 ebenfalls erfüllt. Nun gelten die unmittelbar nach der Definition von $M(f(x))$ in 3 sowie in Anmerkung 13) gemachten Eindeutigkeits-Überlegungen, sobald 1–5 erfüllt sind. Also ist $M'(f(x)) = M(f(x))$, d.h. es gilt auch 7.

Also ist $M(f(x))$ tatsächlich ein Mittel stetiger Funktionen, u.zw. das einzige.

5. Die in den Gleichungen *, ** zum Ausdruck kommenden Resultate sind noch verschiedener Verschärfungen fähig, von denen die folgende erwähnt sei: Wenn wir die Betrachtungen von 2 nicht mit einer stetigen Funktion $f(x)$, sondern mit mehreren $f'(x), \dots, f^k(x)$ auf einmal vornehmen und an Stelle von $S(f(x))$ $\sum_{i=1}^k S(f^i(x))$ minimisieren, so erhalten wir * (bzw. **) für alle $f^i(x)$ auf einmal. D.h. zu irgendwelchen stetigen $f'(x), \dots, f^k(x)$ und $\varepsilon > 0$ können $N, a_1, \dots, a_N$ so gewählt werden, daß * für alle $f^i(x)$ auf einmal, mit denselben $N, a_1, \dots, a_N$ gilt. Dieser Umstand kann übrigens auch zum Beweise von 2 für $M(f(x))$ verwendet werden.

Ferner kann er mit Hinblick auf die Separabilität der Funktionenmenge der stetigen $f(x)$ unter geeigneter Anwendung des Diagonalverfahrens dazu verwendet werden, eine Folge von Elementensystemen $a_1^{(N)}, \dots, a_N^{(N)}, N=1, 2, \dots$ aus G zu konstruieren, so daß für jedes stetige $f(x)$

gleichmäßig in x gegen $M(f(x))$ konvergiert. (Analog $b_1^{(N)}, \dots,$

$b_N^{(N)}, N=1, 2, \dots$, mit $\frac{\sum_{n=1}^N f(b_n^{(N)} \cdot x)}{N}$ für $N \rightarrow \infty$ gleichmäßig in

x gegen $M(f(x))$ konvergent.) Die Punktsysteme $a_1^{(N)}, \dots, a_N^{(N)}$ liegen also sozusagen gleichmäßig dicht in G .

(Eingegangen den 7. September 1933.)

(Zusatz während der Korrektur:)

Inzwischen gelang es, das hier definierte „Integralmittel“ auf nicht-kompakte Gruppen auszudehnen, wo es indessen nicht ein Analogon des Haarschen Integrals, sondern des H. Bohrschen Integralmittels fastperiodischer Funktionen wird. Dadurch wird der Aufbau einer allgemeinen Theorie fastperiodischer Funktionen in beliebigen (separablen) topologischen Gruppen ermöglicht, sowie eine vollständige Theorie ihrer orthogonalen Darstellungen. Die Ausführung erscheint demnächst in den *Annals of Mathematics*.

ALMOST PERIODIC FUNCTIONS IN A GROUP I

INTRODUCTION

1. The object of the present paper is to extend H. Bohr's famous theory of almost periodic functions [4, I]† to arbitrary groups, and to show that it gives just the maximum range over which the fundamental results of Frobenius-Schur representation theory [21; 22; 30] and its extensions by Peter and Weyl [32] hold. We shall see in particular that all bounded linear representations of a group are equivalent to unitary representations and belong to this class. Another point of importance is that we free ourselves completely from all topological assumptions (such as continuity, etc.) by the use of a definition of almost periodicity due to Bochner [2]. Thus we find that the general theory, which applies to every group $\mathfrak{G}$ whatsoever, is completely free from topological assumptions, but all of its results (for example, all series expansions) have a property of closure; if applied to functions which are continuous in a certain topology, they will lead only to functions of the same kind. It is remarkable that we find in the classical case of Bohr new almost periodic functions in addition to the known ones; even the elementary functions $f(a) = e^{2\pi\lambda a}$ can be generalized (this connects with results of Ursell [28]). On the other hand, in some groups (for example in all semi-simple Lie groups) almost periodicity automatically implies continuity (this will be proved with the aid of a theorem of van der Waerden [29]).

2. The principal difficulty in building up a general theory of almost periodic functions lies in finding a generalization of the Bohr integral mean

$$\lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T f(x) dx$$

if the real numbers x and T are replaced by the elements of an arbitrary group $\mathfrak{G}$ which need not be even topological; also, the function $f(x)$ may be discontinuous. We meet this difficulty by finding an entirely new definition (cf. Definitions 4 and 5) which may be proved to be fit for the role of a "mean" under all conditions. The direct discussion of our mean is very simple and is given in Part I.

† The numbers in brackets refer to the bibliography at the end of the paper.

This mean is an extension of an integral in compact groups previously defined by the author [19]. It is defined by means entirely different from those employed in Haar's integral [11] with which it coincides for compact groups, but from which it differs widely for non-compact groups, first, because for such groups it is an integral-mean and not an integral; second, because it is free from topological restrictions, while Haar's integral applies only to locally compact and separable groups; third, because it is defined for almost periodic functions while Haar's integral is defined for measurable functions, and in general neither of these two classes contains the other.

3. The content of Parts I–V is as follows: Part I gives our general theory of the mean. Part II applies this theory (by using the powerful method of Weyl [31]) to prove the fundamental theorems of the Bohr theory, Parseval's formula and the approximation theorem. As we have to combine the devices contained in two papers of Weyl [30; 31] we find it advisable to give the proofs in full, even though the repetition is often almost literal. Part III repeats the main results of the Frobenius-Schur and Peter-Weyl theory of representations, and connects them with the theory of almost periodic functions. It provides a basis for the statement that the present general theory of almost periodic functions is the widest range over which this theory of representations holds without any loss of strength. Part IV connects our theory with topological and other restrictive conditions. By investigating the details of eight examples we illustrate the principal types of combinations of these notions which are likely to occur. Finally, we discuss the question as to how many almost periodic functions exist in a given group. Part V is entirely devoted to the proof that the maximal amount exists in Abelian groups (subject, however, to certain topological restrictions). Here the integral of Haar is used in combination with certain theorems of the author on operators and functions of operators [17]. The extension of some results of Haar on countably infinite Abelian groups [10] is of great importance for these investigations.

4. It is probable that most of the further developments of the Bohr theory will also apply to our general theory. Among these developments are finer convergence theorems, summability theorems, and Stepanoff's generalizations (where some topological restrictions will be necessary, as the Haar integral must be applied). In this connection it may be of interest to point out a needed generalization of an important notion of the Bohr theory, namely, the fact that the product of two elementary almost periodic functions is a function of the same kind: $e^{2\pi\lambda a i} e^{2\pi\mu a i} = e^{2\pi(\lambda+\mu) a i}$. This is unchanged for Abelian groups and leads to the important character-group; but in non-Abelian groups the corresponding situation is that the direct product of two irreducible representations (the elements $D_{\rho\sigma}(a; \mathbb{C})$ of which are the analogues

of $e^{2\pi\lambda a i}$, cf. Definitions 11 and 12, and Theorems 24 and 28) is a sum of a finite number of irreducible representations, that is, there is the so-called composition formula

$$(*) \quad D_{\rho\sigma}(a; \mathfrak{G}) D_{\tau\nu}(a; \mathfrak{D}) = \sum_{\xi, \eta} \Gamma_{\xi, \eta}(\rho, \sigma; \mathfrak{G} | \tau, \nu; \mathfrak{D}) D_{\xi\eta}(a; \mathfrak{G}).$$

Another important notion in Bohr's theory is the independence of the expansion functions $e^{2\pi\lambda a i}$, $e^{2\pi\mu a i}$, $\dots$ (that is, the linear independence of their exponents with integral coefficients), since almost periodic functions with such expansions possess particularly simple convergence properties. The corresponding requirement in our general theory is probably that the right-hand member of (*) should contain no term originating from the representation $D(a; \mathfrak{G}) \equiv 1$ if the left-hand member is any product of powers of $D(a; \mathfrak{G})$, $D(a; \mathfrak{D})$, $\dots$.

I. EXISTENCE OF THE MEAN, GENERAL PROPERTIES

5. Let $\mathfrak{G}$ be a group, that is, a set in which the operations ab and a^{-1} are defined and satisfy the group postulates. While $\mathfrak{G}$ may be topological* this property is not needed in Parts I–III and we do not yet make this assumption concerning $\mathfrak{G}$. Elements of $\mathfrak{G}$ will be denoted by $a, b, c, x, y, z, \dots$, real or complex numbers by $m, n, u, v, \alpha, \beta, \xi, \eta, \dots$, and functions defined in $\mathfrak{G}$ with complex numbers as values by $f(x), g(x), \dots$.

For such functions $f(x)$ and $g(x)$ we define distance† by

$$D(f, g) = \text{l.u.b.}_x |f(x) - g(x)|.$$

A set $\mathfrak{M}$ of such functions is called conditionally compact (c.c.) if every sequence $f_1, f_2, \dots$ extracted from it contains a subsequence $f_{n_1}, f_{n_2}, \dots$ such that $D(f_{n_\mu}, f_{n_\nu}) \rightarrow 0$ as $\mu, \nu \rightarrow \infty$ (that is, a "fundamental" subsequence [13, p. 107]); this means that there exists a function f (not necessarily belonging to $\mathfrak{M}$) such that $D(f_{n_\mu}, f) \rightarrow 0$ as $\mu \rightarrow \infty$.

We now extend Bohr's notion of almost periodic functions [4; 2, §5] to all $f(x)$ in $\mathfrak{G}$, but we prefer to generalize the definition given by S. Bochner [2], as it allows us to rid ourselves completely of topological conditions on $f(x)$ (continuity, etc.).

* That is, a topological set in the sense of Hausdorff [13, pp. 226–230]. One may take his topological system based on the notion of a neighborhood by means of Axioms 1, 2, 3 (or A, B, C) and one of the "separation" Axioms 4–8, such as 5. Furthermore, certain continuity assumptions have to be made concerning ab and a^{-1} . In Parts I–III we shall need no topology at all, in Part IV we must assume that ab is continuous in a for fixed b and in b for fixed a , and in Part V we must assume that ab is continuous in (a, b) and that a^{-1} is continuous in a .

† We shall consider only bounded functions. l.u.b._x denotes the least upper bound for all x 's in $\mathfrak{G}$.

DEFINITION 1. A function $f(x)$ in $\mathfrak{G}$ (with complex values) is called *right almost periodic (r.a.p.)* if the set R_f of all functions $f(xa)$ (x is the variable, a is a parameter running over $\mathfrak{G}$) is c.c.; it is called *left almost periodic (l.a.p.)* if the set L_f of all functions $f(ax)$ is c.c.; it is called *almost periodic (a.p.)* if it is r.a.p. and l.a.p.

The equivalence of this definition to the obvious generalization of the Bohr definition is shown in the usual way if $f(x)$ is continuous; similarly, the uniform continuity of $f(x)$ follows in this case. But as we do not wish now to assume any topology in $\mathfrak{G}$, we shall not go into the details of this matter. On the other hand, the following theorems are of major importance:

THEOREM 1. Each of the three notions r.a.p., l.a.p. and a.p. is invariant under the following operations: $f(xa)$, $f(ax)$, $\overline{f(x)}$, $\alpha f(x)$ (α any complex number), $f(x) \pm g(x)$, $f(x)g(x)$, and the operation of passing from $f_1(x), f_2(x), \dots$ to $f(x)$ if $f_n(x)$ converges uniformly to $f(x)$ as $n \rightarrow \infty$. Passing from $f(x)$ to $f(x^{-1})$ interchanges r.a.p. and l.a.p. and leaves a.p. invariant.

The statement concerning $f(x^{-1})$ is obvious. In the other cases we need to consider only r.a.p., as l.a.p. results, for example, by replacing ab by ba when defining $\mathfrak{G}$, and a.p. results by combining r.a.p. and l.a.p. That $f(xa)$ is r.a.p. is seen by replacing $a_1, a_2, \dots$ (Definition 1) by $a_1a, a_2a, \dots$; that $f(ax)$ is r.a.p. results from replacing x by ax ; the situation concerning $\overline{f(x)}$ and $\alpha f(x)$ is obvious; the r.a.p. of $f(x) \pm g(x)$ and of $f(x)g(x)$ is proved by applying Definition 1 first to $f(x)$ and $a_1, a_2, \dots$, and then to $g(x)$ and the subsequence which has been selected. An obvious and simple application of the diagonal process shows the invariance of r.a.p. under the operation of passing from $f_n(x)$ to $f(x)$.

THEOREM 2. Every r.a.p. or l.a.p. function $f(x)$ is bounded.

Again it is sufficient to consider r.a.p. If $f(x)$ were not bounded, we could select a sequence $a_1, a_2, \dots$ such that $|f(a_n)| \rightarrow \infty$ as $n \rightarrow \infty$, and then no subsequence of $f(xa_1), f(xa_2), \dots$ could have a finite limit at $x = 1$.

DEFINITION 2. If $\mathfrak{M}$ is a set of functions in $\mathfrak{G}$, we call the set of all functions $\alpha_1 f_1(x) + \dots + \alpha_n f_n(x)$ ($n = 1, 2, \dots$; $\alpha_1, \dots, \alpha_n$ non-negative real numbers such that $\alpha_1 + \dots + \alpha_n = 1$; $f_1, \dots, f_n$ any elements of $\mathfrak{M}$) the *convex* of $\mathfrak{M}$ and denote it by $\text{Co}(\mathfrak{M})$.

6. We prove

THEOREM 3. If either of the sets $\mathfrak{M}$ and $\text{Co}(\mathfrak{M})$ is c.c., the other is also c.c.

If $\text{Co}(\mathfrak{M})$ is c.c., its subset $\mathfrak{M}$ is c.c. Conversely, suppose that $\mathfrak{M}$ is c.c. The c.c. property of a set $\mathfrak{M}$ is equivalent to the following condition: for every

$\epsilon > 0$ there exists a finite number of functions $\bar{f}_1, \dots, \bar{f}_m$ of $\mathfrak{N}$ ($m = m(\epsilon)$) such that, for each $f \in \mathfrak{N}$, some $D(f, \bar{f}_\mu) \leq \epsilon$, $\mu = 1, \dots, m$ [13, pp. 108–109]. Now if an $\epsilon > 0$ is given, choose the functions $\bar{f}_1, \dots, \bar{f}_m$ for $\mathfrak{M}$ and ϵ , put

$$\max_{\mu} \text{l.u.b.} |\bar{f}_\mu(x)| = C,$$

and select an integer $N \geq Cm\epsilon^{-1}$. Then $\beta_1\bar{f}_1 + \dots + \beta_m\bar{f}_m$ (where $\beta_1, \dots, \beta_m$ are non-negative rational numbers with denominators N such that $\beta_1 + \dots + \beta_m = 1$) can be written as a finite sequence $\bar{g}_1, \dots, \bar{g}_M$ and have the property described above for $\text{Co}(\mathfrak{M})$ and 2ϵ .

DEFINITION 3. If $f(x)$ is a real bounded function in $\mathfrak{G}$, we call

$$\text{l.u.b.}_{x, y} |f(x) - f(y)| \quad (x \text{ and } y \text{ vary independently over } \mathfrak{G})$$

the oscillation of $f(x)$ and denote it by $\text{Osc}_x f(x)$. If $f(x)$ is not a constant, $\text{Osc}_x f(x) > 0$; otherwise, $\text{Osc}_x f(x)$ is zero.

THEOREM 4. For every real $g \in \text{Co} R_f$, we have $\text{Osc}_x g(x) \leq \text{Osc}_x f(x)$. If the relation $\text{Osc}_x g(x) < \text{Osc}_x f(x)$ never occurs, and if $f(x)$ is l.a.p., $f(x)$ is necessarily a constant.

The first statement is obvious. Suppose that the assumptions of the second statement are valid. Let $a'_1, \dots, a'_n$ be any elements of $\mathfrak{G}$; then

$$\frac{f(xa'_1) + \dots + f(xa'_n)}{n} \in \text{Co } R_f,$$

and thus

$$\text{Osc}_x \frac{f(xa'_1) + \dots + f(xa'_n)}{n} = \text{Osc}_x f(x).$$

This implies that

$$\text{l.u.b.}_x \frac{f(xa'_1) + \dots + f(xa'_n)}{n} = \text{l.u.b.}_x f(x).$$

Put $\text{l.u.b.}_x f(x) = C$; then for every $\epsilon > 0$ there exists an x' such that

$$\frac{f(x'a'_1) + \dots + f(x'a'_n)}{n} \geq C - \epsilon.$$

As all $f(x'a'_i) \leq C$, they all must also be $\geq C - n\epsilon$.

Now choose a $\delta > 0$ and find a finite number of elements of L_f such that each element of L_f has a distance $\leq \delta$ from one of them (cf. the proof of The-

orem 3). That is, find a finite number of elements $a_1, \dots, a_n$ of $\mathfrak{G}$ such that for every a of $\mathfrak{G}$ there exists a $\mu = 1, 2, \dots, n$ for which $|f(a_\mu x) - f(ax)| \leq \delta$ identically. Now choose a b of $\mathfrak{G}$ and repeat the argument just described in the case where $\epsilon = \delta/n$, $a'_1 = a_1^{-1}b, \dots, a'_n = a_n^{-1}b$. Thus an x' exists for which all $f(x'a_\nu^{-1}b) \geq C - \delta$, $\nu = 1, 2, \dots, n$, and therefore, for a properly chosen $\mu = 1, 2, \dots, n$, all $f(a_\mu a_\nu^{-1}b) \geq C - 2\delta$. If $\nu = \mu$, then $f(b) \geq C - 2\delta$.

On the other hand, $f(b) \leq C$ and, as δ was arbitrary, it follows that $f(b) = C$. Finally, $f(x)$ is constant since b was arbitrary. This completes the proof of the theorem.

THEOREM 5. *If $f(x)$ is a.p., there exists a constant A toward which a certain sequence extracted from $\text{Co}R_f$ converges uniformly.*

Since $f(x)$ is r.a.p., R_f and $\text{Co}R_f$ are c.c. Denote real and imaginary parts by $\Re$ and $\Im$ respectively, consider the non-negative numbers $\text{Osc}_x \Re g(x) + \text{Osc}_x \Im g(x)$, $g \in \text{Co}R_f$ and call their greatest lower bound ω . We can extract a sequence $g_1(x), g_2(x), \dots$ from $\text{Co}R_f$ such that $\text{Osc}_x \Re g_n(x) + \text{Osc}_x \Im g_n(x) \rightarrow \omega$ as $n \rightarrow \infty$, and from this a subsequence $g_{n_1}(x), g_{n_2}(x), \dots$, which converges uniformly to a function $g(x)$. Hence $\text{Osc}_x \Re g(x) + \text{Osc}_x \Im g(x) = \omega$. It is obvious that, $f(x)$ being l.a.p., every element $f(xa)$ of R_f is l.a.p. Therefore every element of $\text{Co}R_f$ is l.a.p., and the uniform limit $g(x)$ as well as the real functions $\Re g(x)$ and $\Im g(x)$ are l.a.p. If we show that $\text{Osc}_x \Re g(x) = \text{Osc}_x \Im g(x) = 0$, we have $\Re g(x) = \text{constant}$, $\Im g(x) = \text{constant}$, that is, $g(x) = \text{constant}$, which proves our statement.

Suppose that $\text{Osc}_x \Re g(x) > 0$. Then Theorem 4 shows that an $h \in \text{Co}R_{\Re g}$ exists such that $\text{Osc}_x h(x) < \text{Osc}_x \Re g(x)$. Here $h(x) = \alpha_1 \Re g(xa_1) + \dots + \alpha_n \Re g(xa_n)$ ($\alpha_1, \dots, \alpha_n$ each ≥ 0 , $\alpha_1 + \dots + \alpha_n = 1$). Putting $k(x) = \alpha_1 g(xa_1) + \dots + \alpha_n g(xa_n)$, we have $h(x) = \Re k(x)$, so that $\text{Osc}_x \Re k(x) < \text{Osc}_x \Re g(x)$. But it is obvious that $\text{Osc}_x \Im k(x) \leq \text{Osc}_x \Im g(x)$. Therefore $\text{Osc}_x \Re k(x) + \text{Osc}_x \Im k(x) < \omega$. Now $g(x)$ can be uniformly approximated by functions $l \in \text{Co}R_f$, that is, $l(x) = \beta_1 f(xb_1) + \dots + \beta_m f(xb_m)$ ($\beta_1, \dots, \beta_m$ each ≥ 0 , $\beta_1 + \dots + \beta_m = 1$). Hence $k(x)$ can be uniformly approximated by functions $q(x) = \alpha_1 \beta_1 f(xa_1 b_1) + \alpha_1 \beta_2 f(xa_1 b_2) + \dots + \alpha_n \beta_m f(xa_n b_m)$, that is, by functions $q \in \text{Co}R_f$. Since $\text{Osc}_x \Re k(x) + \text{Osc}_x \Im k(x) < \omega$, the relation that $\text{Osc}_x \Re q(x) + \text{Osc}_x \Im q(x) < \omega$ results. This contradicts the definition of ω . Similarly $\text{Osc}_x \Im g(x) > 0$ is disproved.

REMARK. *If a finite number of a.p. functions $f_1(x), \dots, f_t(x)$ are given, it is possible to find a set of constants $A_1, \dots, A_t$, toward which t sequences extracted from $\text{Co}R_{f_1}, \dots, \text{Co}R_{f_t}$ respectively, with the same $\alpha_1, \dots, \alpha_n, a_1, \dots, a_n$, converge uniformly (that is, sequences of the form $\alpha_1^{(\nu)} f_1(xa_1^{(\nu)})$*

$+\cdots+\alpha_{n_r}^{(r)}f_1(xa_{n_r}^{(r)}), \cdots, \alpha_{1t}^{(r)}f(xa_1^{(r)})+\cdots+\alpha_{n_r}^{(r)}f_t(xa_{n_r}^{(r)}),$ where $\nu \rightarrow \infty, \alpha_1^{(\nu)} \geq 0, \cdots, \alpha_{n_r}^{(\nu)} \geq 0,$ and $\alpha_1^{(\nu)} + \cdots + \alpha_{n_r}^{(\nu)} = 1$.

The argument which proved Theorem 5 may be repeated here if we use $\text{Osc}_x \Re f_1(x) + \text{Osc}_x \Im f_1(x) + \cdots + \text{Osc}_x \Re f_t(x) + \text{Osc}_x \Im f_t(x)$ instead of $\text{Osc}_x \Re f(x) + \text{Osc}_x \Im f(x)$.

DEFINITION 4. A real number A which may be uniformly approximated by functions from $\text{Co}R_f$ or $\text{Co}L_f$, that is, a number A such that, for every $\epsilon > 0$, there exists a number $n=1, 2, \cdots$, numbers $\alpha_1, \cdots, \alpha_n$ each ≥ 0 with $\alpha_1 + \cdots + \alpha_n = 1$, and elements $a_1, \cdots, a_n$ of $\mathfrak{G}$ such that the condition $|\alpha_1 f(xa_1) + \cdots + \alpha_n f(xa_n) - A| \leq \epsilon$ or $|\alpha_1 f(a_1x) + \cdots + \alpha_n f(a_nx) - A| \leq \epsilon$ holds throughout $\mathfrak{G}$, is called a right-mean or a left-mean of $f(x)$ respectively.

THEOREM 6. If $f(x)$ is a.p., it has exactly one right-mean, exactly one left-mean, and these means are equal.

The existence of a right-mean has been proved by Theorem 5. If we change the multiplication law ab in $\mathfrak{G}$ to ba , all notions remain unchanged except for the interchange of "right" and "left." Thus a left-mean must exist.

Now let A be a right-mean, let B be a left-mean, and let ϵ be > 0 . Choose $\alpha_1, \cdots, \alpha_n, \beta_1, \cdots, \beta_m, a_1, \cdots, a_n, b_1, \cdots, b_m$ such that

$$\begin{aligned} |\alpha_1 f(xa_1) + \cdots + \alpha_n f(xa_n) - A| &\leq \epsilon, \\ |\beta_1 f(b_1x) + \cdots + \beta_m f(b_mx) - B| &\leq \epsilon. \end{aligned}$$

If we replace x in the first equation by $b_1x, \cdots, b_mx$ in succession, and add, we obtain

$$|\alpha_1\beta_1 f(b_1xa_1) + \alpha_1\beta_2 f(b_2xa_1) + \cdots + \alpha_n\beta_m f(b_mxa_n) - A| \leq \epsilon.$$

Similarly, if we replace x in the second equation by $xa_1, \cdots, xa_n$ in succession, we obtain

$$|\alpha_1\beta_1 f(b_1xa_1) + \alpha_1\beta_2 f(b_2xa_1) + \cdots + \alpha_n\beta_m f(b_mxa_n) - B| \leq \epsilon.$$

Therefore $|A - B| \leq 2\epsilon$ and, as ϵ may be arbitrarily small, $A = B$.

DEFINITION 5. If $f(x)$ is a.p., we call the common value of its uniquely determined right- and left-means the mean of $f(x)$, and denote it by $M_x f(x)$.*

We now state the most important properties of the mean.

* Definitions 3-5 and the argument of Theorems 3-6 are in very close analogy to the author's construction of the Haar-Lebesgue measure in compact groups [19]. It is noteworthy that for non-compact groups, where Haar proved by his method the existence of an integral [11], our method leads to an integral-mean.

THEOREM 7. *If $f(x)$ and $g(x)$ are a.p. functions, all the functions $f(xa)$, $f(ax)$, $f(x^{-1})$, $\overline{f(x)}$, $\alpha f(x)$, $f(x) \pm g(x)$ (α a complex number, a an element of $\mathfrak{G}$) are a.p. (cf. Theorem 1). Furthermore, we have the following:*

- (1) $M_x[\alpha f(x)] = \alpha M_x f(x)$.
- (2) $M_x[f(x) \pm g(x)] = M_x f(x) \pm M_x g(x)$.
- (3) $M_x 1 = 1$.
- (4) *If $f(x)$ is real and ≥ 0 throughout $\mathfrak{G}$, then $M_x f(x) \geq 0$; and if, in addition, $f(x) \not\equiv 0$, then $M_x f(x) > 0$.*
- (5) $|M_x[f(x)]| \leq M_x[|f(x)|]$.
- (6) $M_x[\overline{f(x)}] = \overline{M_x[f(x)]}$.
- (7) $M_x f(xa) = M_x f(x)$.
- (8) $M_x f(ax) = M_x f(x)$.
- (9) $M_x f(x^{-1}) = M_x f(x)$.

The equations (1), (3), (5), (6) and the first half of (4) are obvious; as every left-mean of $f(x)$ is a left-mean of $f(xa)$ and as every right-mean of $f(x)$ is a right-mean of $f(ax)$, (7) and (8) are valid; as every right-mean of $f(x)$ is a left-mean of $f(x^{-1})$, (9) is true. Thus, only (2) and the second half of (4) remain unproved.

In order to prove (2), put $M_x f(x) = A$, $M_x g(x) = B$, let ϵ be > 0 , and choose $\alpha_1, \dots, \alpha_n$ ($\alpha_1, \dots, \alpha_n$ each ≥ 0 , $\alpha_1 + \dots + \alpha_n = 1$) and $a_1, \dots, a_n$ such that

$$|\alpha_1 f(xa_1) + \dots + \alpha_n f(xa_n) - A| \leq \epsilon.$$

Now $\alpha_1 g(xa_1) + \dots + \alpha_n g(xa_n)$ obviously has the same left-mean as $g(x)$, i.e., B . Therefore we can choose $\beta_1, \dots, \beta_m$ ($\beta_1, \dots, \beta_m$ each ≥ 0 , $\beta_1 + \dots + \beta_m = 1$) and $b_1, \dots, b_m$ such that

$$\alpha_1 \beta_1 g(xb_1 a_1) + \alpha_1 \beta_2 g(xb_2 a_1) + \dots + \alpha_n \beta_m g(xb_m a_n) - B \leq \epsilon.$$

If we replace x in the first inequality by $xb_1, \dots, xb_m$ in succession, and add, we obtain

$$|\alpha_1 \beta_1 f(xb_1 a_1) + \alpha_1 \beta_2 f(xb_2 a_1) + \dots + \alpha_n \beta_m f(xb_m a_n) - A| \leq \epsilon.$$

Denote nm by p ; $\alpha_1 \beta_1, \alpha_1 \beta_2, \dots, \alpha_n \beta_m$ by $\gamma_1, \dots, \gamma_p$ ($\gamma_1, \dots, \gamma_p$ each ≥ 0 , $\gamma_1 + \dots + \gamma_p = 1$); $b_1 a_1, b_2 a_1, \dots, b_m a_n$ by $c_1, \dots, c_p$; we get, by adding and subtracting our inequalities,

$$|\gamma_1(f(xc_1) \pm g(xc_1)) + \dots + \gamma_p(f(xc_p) \pm g(xc_p)) - (A \pm B)| \leq 2\epsilon.$$

As ϵ may be arbitrarily small, this shows that $M_x[f(x) \pm g(x)] = A \pm B$. (An-

other way to prove (2) would be to apply the Remark following Theorem 5 to $f(x)$ and $g(x)$.)

In order to prove the second half of (4), assume $f(x) \geq 0$ everywhere and $f(x_0) > 0$ for one particular x_0 . For any $\epsilon > 0$ a finite number of elements of R_f exist such that each element of R_f has a distance $\leq \epsilon$ from one of them (cf. the proof of Theorem 3). Hence there is a finite number of the elements $a_1, \dots, a_n$ such that, for every a , there exists a $\mu = 1, 2, \dots, n$ such that $|f(xa_\mu) - f(xa)| \leq \epsilon$ identically. Now take $\epsilon = f(x_0)/2$. The substitution $x = x_0 a_\mu^{-1}$ shows that $f(x_0 a_\mu^{-1} a) \geq f(x_0)/2$. Hence, for each a it follows that $f(x_0 a_\nu^{-1} a) \geq 0$ for every $\nu = 1, \dots, n$, but that $f(x_0 a_\nu^{-1} a) \geq f(x_0)/2$ for at least one ν . Thus $f(x_0 a_1^{-1} a) + \dots + f(x_0 a_n^{-1} a) \geq f(x_0)/2$, that is, the function $g(y) = f(x_0 a_1^{-1} y) + \dots + f(x_0 a_n^{-1} y) - f(x_0)/2$ is always ≥ 0 . Hence the first half of (4) leads to the result that $M_y g(y) \geq 0$, (2), (3), (6), (7), and (8) show that $M_y g(y) = n M_y f(y) - f(x_0)/2$, and it follows that

$$M_y f(y) \geq \frac{f(x_0)}{2n} > 0.$$

THEOREM 8. *The formal properties (1)–(9) determine $M_x f(x)$ uniquely; in fact, (1)–(3), the first half of (4), and (7) or (8) are sufficient.*

It is sufficient to consider (1)–(3), the first half of (4), and (7), as (8) may be obtained by replacing ab in $\mathfrak{G}$ by ba . So assume that a functional $M'_x f(x)$, defined for all a.p. $f(x)$ and satisfying (1)–(3), the first half of (4), and (7), is given.

For every $\epsilon > 0$ we can choose $\alpha_1, \dots, \alpha_n, a_1, \dots, a_n$ ($\alpha_1, \dots, \alpha_n$ each ≥ 0 , $\alpha_1 + \dots + \alpha_n = 1$) such that $|\alpha_1 f(xa_1) + \dots + \alpha_n f(xa_n) - M'_x f(x)| \leq \epsilon$, or if $f(x)$ is real,

$$M'_x f(x) - \epsilon \leq \alpha_1 f(xa_1) + \dots + \alpha_n f(xa_n) \leq M'_x f(x) + \epsilon.$$

Then (1)–(3), the first half of (4), and (7) show that $M'_x f(x) - \epsilon \leq M'_x f(x) \leq M'_x f(x) + \epsilon$, and as ϵ was arbitrary, $M'_x f(x) = M_x f(x)$. Property (1) with $\alpha = i$ shows that this holds also for pure imaginary $f(x)$, and property (2) shows that it holds for every $f(x)$.

Theorems 6–8 show that, for a.p. functions $f(x)$, there is exactly one way to define a notion $M_x f(x)$ possessing the essential formal properties of a mean. Our $M_x f(x)$ is the equivalent of the well known integral mean

$$\lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T f(x) dx$$

in Bohr's theory, when $\mathfrak{G}$ is the addition group of all real numbers. But even

in this case the form of our definition is essentially different from the usual one (for example, it does not use continuity), and it gives a new approach to the problem.

REMARK. The notion of the mean can be modified in the following manner. Consider the doubled group $\mathfrak{G}\mathfrak{G}'$, that is, the set of all pairs $[a, a']$. This set is a group by virtue of the definitions $[a, a'] [b, b'] = [ab, b'a']$ and $[a, a']^{-1} = [a^{-1}, a'^{-1}]$. (This is similar to the construction in the next paragraph except that here we use $b'a'$ while there we use $a'b'$.) The argument in the proof of Theorem 9 below shows that if $f(x)$ is a.p. in $\mathfrak{G}$, then $f_0([x, x']) = f(xx')$ is a.p. in $\mathfrak{G}\mathfrak{G}'$. By Theorem 5 it then follows that there exists a constant A such that, for every $\epsilon > 0$, there exists a number $n = 1, 2, \dots$, numbers $\alpha_1, \dots, \alpha_n$ each ≥ 0 with $\alpha_1 + \dots + \alpha_n = 1$, and elements $a_1, \dots, a_n$ and $b_1, \dots, b_n$ of $\mathfrak{G}$ such that the condition

$$| \alpha_1 f_0([x, y][a_1, b_1]) + \dots + \alpha_n f_0([x, y][a_n, b_n]) - A | \leq \epsilon$$

holds for all x and y in $\mathfrak{G}$. If we write $c_1 = a_1 b_1, \dots, c_n = a_n b_n$, this condition assumes the form

$$| \alpha_1 f(x c_1 y) + \dots + \alpha_n f(x c_n y) - A | \leq \epsilon.$$

This mean is even easier to handle than our right- and left-means (which are special cases of it). This is due to the following fact: choose two arbitrary sets of numbers $\beta_1, \dots, \beta_k$ and $\gamma_1, \dots, \gamma_l$, each ≥ 0 , with $\beta_1 + \dots + \beta_k = 1$ and $\gamma_1 + \dots + \gamma_l = 1$, and two arbitrary sets of elements $a_1, \dots, a_k$ and $b_1, \dots, b_l$ of $\mathfrak{G}$. In our last inequality replace x and y by $x a_\kappa$ and $b_\lambda y$, multiply by $\beta_\kappa \gamma_\lambda$, and add over all $\kappa = 1, \dots, k$; $\lambda = 1, \dots, l$. Then we obtain an inequality of the same type except that there are knl terms instead of n terms and $\beta_\kappa \gamma_\lambda$ and $a_\kappa c_\lambda b_\lambda$ appear in place of α_ν and a_ν . This shows that the conditions

$$\begin{aligned} | \alpha_1 f(x c_1 y) + \dots + \alpha_m f(x c_m y) - A | &\leq \epsilon, \\ | \alpha'_1 g(x c'_1 y) + \dots + \alpha'_n g(x c'_n y) - B | &\leq \epsilon \end{aligned}$$

imply the conditions

$$\begin{aligned} | \alpha''_1 f(x c''_1 y) + \dots + \alpha''_{mn} f(x c''_{mn} y) - A | &\leq \epsilon, \\ | \alpha'_1 g(x c'_1 y) + \dots + \alpha''_{mn} g(x c''_{mn} y) - B | &\leq \epsilon \end{aligned}$$

if α''_ρ and c''_ρ are $\alpha_\mu \alpha'_\nu$ and $c_\mu c'_\nu$ (in some order). This gives the uniqueness, the extension to complex $f(x)$, and the additivity of our new mean at once. Of course this mean coincides with our former means.

7. The applications to be made in the next chapter necessitate our proving some facts concerning double means. We therefore pass to this subject.

The group $\mathfrak{G}$ can be "doubled," that is, we can consider the set $\mathfrak{G}\mathfrak{G}$ of all pairs $[a, a']$, which by the definitions $[a, a'] [b, b'] = [ab, a'b']$, $[a, a']^{-1} = [a^{-1}, a'^{-1}]$ becomes a group, and we will denote functions in it by $f(x, x')$ instead of by $f([x, x'])$. All our notions apply to $\mathfrak{G}\mathfrak{G}$: we have a.p. functions $f(x, x')$ in $\mathfrak{G}\mathfrak{G}$, and a mean $M_{x, x'} f(x, x')$.

THEOREM 9. *If $f(x)$ is a.p. in $\mathfrak{G}$, the eight functions $f(xx')$, $f(x'x)$, $f(xx'^{-1})$, $\dots$, $f(x'^{-1}x^{-1})$ are all a.p. in $\mathfrak{G}\mathfrak{G}$.*

Interchange of x and x' , of ab in $\mathfrak{G}$ with ba , and of $f(x)$ with $f(x^{-1})$ reduces our task to discussing $f(xx')$ and $f(xx'^{-1})$ alone. Their a.p. character in $\mathfrak{G}\mathfrak{G}$ means that the sets of functions $f(axa'x')$, $f(xax'a')$, $f(axx'^{-1}a'^{-1})$, $f(xaa'^{-1}x'^{-1})$ in $\mathfrak{G}\mathfrak{G}$ are c.c. or else that the sets of functions of one or of two variables, $f(axby)$, $f(xayb)$, $f(axb)$, $f(xay)$ in $\mathfrak{G}$ are c.c. The third case arises from the first by setting $y=1$, the fourth from the second by setting $b=1$, and the second from the first by interchanging a and x with b and y , and ab in $\mathfrak{G}$ with ba . So we need to discuss only $f(axby)$.

Choose an $\epsilon > 0$. As $f(x)$ is r.a.p., there is a finite number of elements $y_1, \dots, y_n$, such that for each y there is a $\nu=1, \dots, n$ for which $|f(zy) - f(zy_\nu)| \leq \epsilon$ identically. As each $f(xy_\nu)$ is r.a.p., the set of all functions $f(xby_\nu)$ is c.c. for every $\nu=1, \dots, n$, and therefore even the set of all "vector-functions" with n components $[f(xby_1), \dots, f(xby_n)]$ is c.c. Therefore a finite number of elements $b_1, \dots, b_m$ exist such that for each b there is a $\mu=1, 2, \dots, m$ for which $|f(xby_\nu) - f(xb_\mu y_\nu)| \leq \epsilon$ for all x and for every $\nu=1, 2, \dots, n$. Our two inequalities together give the result that

$$|f(xby) - f(xb_\mu y)| \leq 3\epsilon.$$

Finally, $f(x)$ is l.a.p., so that there exists a finite set of elements $a_1, \dots, a_l$ such that for every a there is a $\lambda=1, \dots, l$ for which $|f(au) - f(a_\lambda u)| \leq \epsilon$ identically. This, together with our last inequality, implies that $|f(axby) - f(a_\lambda x b_\mu y)| \leq 5\epsilon$ identically.

As this holds for every $\epsilon > 0$, the c.c. character of the set of functions $f(axby)$ is proved [13, pp. 108-109].

THEOREM 10. *If $f(x, x')$ is a.p. in $\mathfrak{G}\mathfrak{G}$, it is also a.p. in $\mathfrak{G}$ as a function of x or as a function of x' . Thus we can form $M_x f(x, x')$ and $M_{x'} f(x, x')$ which are a.p. in $\mathfrak{G}$ as a function of x' and as a function of x , respectively. Thus we can form $M_{x'} [M_x f(x, x')]$ and $M_x [M_{x'} f(x, x')]$. These expressions are both equal to $M_{xx'} f(x, x')$.*

The first statement is obvious. In the second and third statements it is sufficient to consider $M_x f(x, x')$ and $M_{x'} [M_x f(x, x')]$, as interchange of x

with x' and of $f(x, x')$ with $f(x', x)$ leads to the rest of the theorem.

Consider a sequence $a_1, a_2, \dots$ of elements of $\mathfrak{G}$. As $f(x, x')$ is r.a.p., the sequence $f(x, x'a_1), f(x, x'a_2), \dots$ contains a uniformly convergent subsequence $f(x, x'a_{n_1}), f(x, x'a_{n_2}), \dots$ such that, for every $\epsilon > 0$ and almost all μ and ν , $|f(x, x'a_{n_\mu}) - f(x, x'a_{n_\nu})| \leq \epsilon$. This implies that $|M_{x'}f(x, x'a_{n_\mu}) - M_{x'}f(x, x'a_{n_\nu})| \leq \epsilon$. Thus the set of functions of x' , $M_{x'}f(x, x'a)$, is c.c. Therefore $M_{x'}f(x, x')$ is r.a.p. and interchange of ab in $\mathfrak{G}$ with ba shows that it is l.a.p. Hence it is a.p.

Now it is obvious that $M'_{xx'}f(x, x') = M_{x'}[M_{x'}f(x, x')]$ has Properties (1)–(4) and (7) enumerated in Theorem 7 if we look at it as an $[x, x']$ -mean. Therefore we may conclude from Theorem 8 that it is $M_{xx'}f(x, x')$.

Theorems 9 and 10 may be extended by iterating them m times to functions of 2^m variables; by choosing $2^m \geq n$ and taking the functions constant in the last $2^m - n$ variables these theorems may be extended to functions of n variables.

II. APPLICATION OF THE METHOD OF WEYL AND E. SCHMIDT • PROOF OF THE FUNDAMENTAL THEOREMS

8. The results of Part I enable us to apply the method of Weyl to the proof of the fundamental theorems of Bohr's theory of a.p. functions (in the addition group of real numbers) and to the discussion of the linear-orthogonal representations of continuous groups.* The present part, II, contains a proof of "Parseval's formula" (equivalent to Theorem 15), which runs exactly along the lines of Weyl's proof. It also contains the proof of the "approximation theorem" (equivalent to Theorem 18) where a different device, due to N. Wiener, has to be used because of the difficulties of constructing in our general case an a.p. function with the required properties (cf. [31, pp. 348–349], and our Theorem 17). The next part, III, contains an interpretation and application of these theorems connecting the theories of a.p. functions and of representations. In this, Weyl's method is of fundamental importance.

DEFINITION 6. If $f(x)$ and $g(x)$ are a.p., we set

$$h(x) = M_y[f(xy^{-1})g(y)] = M_y[f(y)g(y^{-1}x)] = f \times g.$$

We observe that the two expressions for $h(x)$ are equal by Theorem 7 and Properties (7) and (9), after making the substitution $y^{-1}x$ for y , and that $h(x)$ is a.p. by Theorems 9 and 10.

* Cf. H. Weyl [31], H. Weyl and F. Peter [32]. The operational methods used there are partly based on the thesis of E. Schmidt [20].

REMARK. $f \times g$ can be uniformly approximated by functions of the form $\gamma_1 f(xc_1) + \dots + \gamma_n f(xc_n)$ ($\gamma_1, \dots, \gamma_n$ are complex numbers), that is, for every $\epsilon > 0$ there exist numbers $\gamma_1, \dots, \gamma_n$ and elements $c_1, \dots, c_n$ such that

$$|f \times g(x) - \gamma_1 f(xc_1) - \dots - \gamma_n f(xc_n)| \leq \epsilon$$

holds throughout $\mathfrak{G}$.

$g(x)$ is a.p. and therefore bounded (Theorem 2); suppose $g(x) \leq C$. Now choose a $\delta > 0$. According to our remark in the proof of Theorem 3, it is possible to find a finite number of elements $b_1, \dots, b_k$ of $\mathfrak{G}$ such that to every x there exists a $\kappa = 1, \dots, k$ for which $|f(xz) - f(b_\kappa z)| \leq \delta$ holds identically. Now consider the a.p. y -functions $f(b_\kappa y^{-1})g(y)$, $\kappa = 1, \dots, k$, and apply to them the Remark following Theorem 5 (with $\epsilon/2$): if $\epsilon > 0$, there exists a set of real numbers $\alpha_1, \dots, \alpha_n$ ($\alpha_1, \dots, \alpha_n$ each ≥ 0 , $\alpha_1 + \dots + \alpha_n = 1$) and a set of elements $a_1, \dots, a_n$ of $\mathfrak{G}$ such that

$$|\alpha_1 f(b_\kappa a_1^{-1} y^{-1})g(ya_1) + \dots + \alpha_n f(b_\kappa a_n^{-1} y^{-1})g(ya_n) - M_y[f(b_\kappa y^{-1})g(y)]| \leq \frac{\epsilon}{2}$$

holds identically for every $\kappa = 1, \dots, k$. For every x there exists a κ such that $|f(xz) - f(b_\kappa z)| \leq \delta$ and therefore such that

$$\begin{aligned} |f(xu^{-1})g(u) - f(b_\kappa u^{-1})g(u)| &\leq C\delta \quad \text{and} \\ |M_u[f(xu^{-1})g(u)] - M_u[f(b_\kappa u^{-1})g(u)]| &\leq C\delta. \end{aligned}$$

Hence we have the result that

$$\begin{aligned} |\alpha_1 f(xa_1^{-1} y^{-1})g(ya_1) + \dots + \alpha_n f(xa_n^{-1} y^{-1})g(ya_n) - M_y[f(xy^{-1})g(y)]| \\ \leq 2C\delta + \frac{\epsilon}{2}. \end{aligned}$$

Our statement is proved if we put $\delta = \epsilon/(4C)$ and $y = 1$, and substitute $\alpha_1 g(a_1), \dots, \alpha_n g(a_n)$ for $\gamma_1, \dots, \gamma_n$ and $a_1^{-1}, \dots, a_n^{-1}$ for $c_1, \dots, c_n$.

THEOREM 11. The "multiplication" $f \times g$ is distributive (linear) in both factors, associative, and if $\mathfrak{G}$ is Abelian, commutative.

The theorem is obvious except for associativity. Our second form for $h = f \times g$ gives

$$\begin{aligned} (f \times g) \times k(x) &= M_z[M_y[f(y)g(y^{-1}z)]k(z^{-1}x)] \\ &= M_z[M_y[f(y)g(y^{-1}z)k(z^{-1}x)]], \\ f \times (g \times k)(x) &= M_y[f(y)M_z[g(y^{-1}z)k(z^{-1}x)]] \\ &= M_y[M_z[f(y)g(y^{-1}z)k(z^{-1}x)]], \end{aligned}$$

and these expressions are equal by Theorems 9 and 10.

DEFINITION 7. If $f(x)$ is a.p., we denote $f \times f \times \cdots \times f$ (n factors) by f^n , and $\overline{f(x^{-1})}$ by $f'(x)$. Furthermore, we define

$$Nf = \{M_x[|f(x)|^2]\}^{1/2}.$$

THEOREM 12. Let $f(x)$ and $g(x)$ be a.p. functions. The following formulas hold:

- (1) If $f \neq 0$, $Nf > 0$.
- (2) $N[\alpha f] = |\alpha| Nf$, $N[f \pm g] \leq Nf + Ng$, $N[f \times g] \leq (Nf)(Ng)$.
- (3) $ff'(1) = f'f(1) = (Nf)^2$.
- (4) $|M_x f(x)| \leq Nf$.
- (5) $|f \times g(x)| \leq (Nf)(Ng)$.
- (6) $M_x[|f(x)||g(x)|] \leq (Nf)(Ng)$.

Statements (1), the first part of (2), and (3) are obvious. The second part of (2), after being squared, means that

$$M_x[|f(x) \pm g(x)|^2] \leq M_x[|f(x)|^2] + M_x[|g(x)|^2] + 2(Nf)(Ng),$$

$$|M_x[\Re(f(x)\overline{g(x)})]| \leq (Nf)(Ng).$$

This obviously follows from (6). The third part of (2) again follows from (5) by squaring and applying M_x . (4) follows from (5) by putting $g(x) \equiv 1$, since $f \times 1(x) = M_y[f(y)]$.

Hence we need to prove only (5) and (6). Since

$$|f(y)g(y^{-1}x)| \leq \frac{1}{2}|f(y)|^2 + \frac{1}{2}|g(y^{-1}x)|^2$$

it follows that

$$\begin{aligned} |M_y[f(y)g(y^{-1}x)]| &\leq \frac{1}{2}M_y[|f(y)|^2] + \frac{1}{2}M_y[|g(y^{-1}x)|^2] \\ &\leq \frac{1}{2}M_y[|f(y)|^2] + \frac{1}{2}M_y[|g(y)|^2], \\ |f \times g(x)| &\leq \frac{1}{2}(Nf)^2 + \frac{1}{2}(Ng)^2. \end{aligned}$$

Starting from

$$|f(x)||g(x)| \leq \frac{1}{2}|f(x)|^2 + \frac{1}{2}|g(x)|^2$$

we obtain similarly

$$M_x[|f(x)||g(x)|] \leq \frac{1}{2}(Nf)^2 + \frac{1}{2}(Ng)^2.$$

If we replace f and g by γf and g/γ (γ real and > 0) we see that $|f \times g(x)|$ and $M_x[|f(x)||g(x)|]$ do not exceed

$$\frac{\gamma^2}{2}(Nf)^2 + \frac{1}{2\gamma^2}(Ng)^2.$$

The greatest lower bound of this expression is $(Nf)(Ng)$. This completes the proof of (5) and (6).

THEOREM 13. *Let $f(x)$ be an a.p. function $\neq 0$. Put*

$$\begin{aligned}\Gamma_n &= N[f \times f' \times \cdots]^2 = N[f' \times f \times \cdots]^2 \quad (n \text{ factors}) \\ &= (f \times f')^n(1) = (f' \times f)^n(1).\end{aligned}$$

Then

$$\Gamma_n > 0, \quad (\Gamma_n)^2 \leq \Gamma_{n-1}\Gamma_{n+1}, \quad \Gamma_{m+n} \leq \Gamma_m\Gamma_n.$$

First we prove that the four expressions above for Γ_n are equal. Indeed the first and third expressions for Γ_n are equal, and so are the second and the fourth, since $(g \times h)' = h' \times g'$. The equality of the third and fourth expressions follows from

$$(f \times f')^n = f \times (f' \times f \times \cdots \times f'), \quad (f' \times f)^n = (f' \times f \times \cdots \times f') \times f,$$

and from

$$f \times g(1) = g \times f(1) = M_y[f(y)g(y^{-1})].$$

By (5) of Theorem 12, $|f \times g(1)| \leq (Nf)(Ng)$. If we replace here f and g by $f \times f' \times \cdots$ with $n-1$ and $n+1$ factors respectively, we obtain $\Gamma_n^2 \leq \Gamma_{n-1}\Gamma_{n+1}$. And if we replace f and g in (2) of Theorem 12 by $f \times f' \times \cdots$ or $f' \times f \times \cdots$ with m and n factors respectively, we obtain $\Gamma_{m+n} \leq \Gamma_m\Gamma_n$. That $\Gamma_n \geq 0$ is obvious; but the condition $\Gamma_n = 0$ would imply that $\Gamma_{n-1} = 0$ (because $\Gamma_{n-1}^2 = \Gamma_{n-2}\Gamma_n$ provided that $n \geq 3$), so that $\Gamma_2 = 0$. This means that $N[f \times f'] = 0$, $f \times f'(x) \equiv 0$, hence $N[f]^2 = f \times f'(1) = 0$ (that is, $\Gamma_1 = 0$) and $f \equiv 0$, contrary to our assumption. Thus we have $\Gamma_n > 0$.

THEOREM 14. *Let $f(x)$ be as before, and define $\Gamma_1, \Gamma_2, \dots$ as before. Then as $n \rightarrow \infty$,*

$$\frac{\Gamma_{n+1}}{\Gamma_n} \rightarrow \gamma, \quad \frac{\Gamma_n}{\gamma^n} \rightarrow \kappa, \quad \frac{(f \times f')^n(x)}{\gamma^n} \rightarrow \phi(x) \quad (\text{uniformly}).$$

Furthermore, $0 < \gamma \leq \Gamma_1$, $1 \leq \kappa$, $\phi(x)$ is a.p., and $\phi' = \phi$, $\phi \times \phi = \phi$, $f \times f' \times \phi = \phi \times f \times f' = \gamma\phi$, $\phi(1) = \kappa$.

The formulas of Theorem 13 imply that

$$0 < \frac{\Gamma_2}{\Gamma_1} \leq \frac{\Gamma_3}{\Gamma_2} \leq \cdots \leq \Gamma_1,$$

and therefore Γ_{n+1}/Γ_n has a limit γ as $n \rightarrow \infty$, $0 < \gamma \leq \Gamma_1$. Furthermore,

$$\frac{\Gamma_{n+1}}{\Gamma_n} \leq \gamma, \quad \frac{\Gamma_n}{\gamma^n} \geq \frac{\Gamma_{n+1}}{\gamma^{n+1}},$$

that is,

$$\frac{\Gamma_1}{\gamma} \geq \frac{\Gamma_2}{\gamma^2} \geq \dots > 0,$$

and therefore Γ_n/γ^n has a limit κ as $n \rightarrow \infty$, $\kappa \geq 0$. Finally we have $\Gamma_{m+n} \leq \Gamma_m \Gamma_n$, that is,

$$\Gamma_n \geq \frac{\Gamma_{m+n}}{\Gamma_m} = \frac{\Gamma_{m+1}}{\Gamma_m} \frac{\Gamma_{m+2}}{\Gamma_{m+1}} \dots \frac{\Gamma_{m+n}}{\Gamma_{m+n-1}} \geq \left(\frac{\Gamma_{m+1}}{\Gamma_m} \right)^n.$$

The limiting process $m \rightarrow \infty$ shows that $\Gamma_n \geq \gamma^n$ and $\Gamma_n/\gamma^n \geq 1$, and then the limiting process $n \rightarrow \infty$ shows that $\kappa \geq 1$.

By (4) and (2) of Theorem 12,

$$\begin{aligned} \left| \frac{(f \times f')^n(x)}{\gamma^n} - \frac{(f \times f')^m(x)}{\gamma^m} \right|^2 &= \left| f \times \left(\frac{(f' \times f)^{n-1}}{\gamma^n} - \frac{(f' \times f)^{m-1}}{\gamma^m} \right) \times f'(x) \right|^2 \\ &\leq (Nf)^2 \left(N \left[\frac{(f' \times f)^{n-1}}{\gamma^n} - \frac{(f' \times f)^{m-1}}{\gamma^m} \right] \right)^2 (Nf')^2 \\ &= \Gamma_1^2 \left(\frac{(f' \times f)^{n-1}}{\gamma^n} - \frac{(f' \times f)^{m-1}}{\gamma^m} \right)^2 (1) \\ &= \Gamma_1^2 \left(\frac{(f' \times f)^{2n-2}(1)}{\gamma^{2n}} - 2 \frac{(f' \times f)^{m+n-2}(1)}{\gamma^{m+n}} + \frac{(f' \times f)^{2m-2}(1)}{\gamma^{2m}} \right) \\ &= \Gamma_1^2 \left(\frac{\Gamma_{2n-2}}{\gamma^{2n}} - 2 \frac{\Gamma_{m+n-2}}{\gamma^{m+n}} + \frac{\Gamma_{2m-2}}{\gamma^{2m}} \right). \end{aligned}$$

As m and $n \rightarrow \infty$ the last expression converges to 0; thus the first expression converges uniformly to 0, that is, as $n \rightarrow \infty$, $(f \times f')^n(x)/\gamma^n$ converges uniformly to a limiting function $\phi(x)$. As the functions $(f \times f')^n(x)/\gamma^n$ are a.p., $\phi(x)$ is also a.p.

The relations

$$\begin{aligned} \left(\frac{(f \times f')^n}{\gamma^n} \right)' &= \frac{(f \times f')^n}{\gamma^n}, & \left(\frac{(f \times f')^n}{\gamma^n} \right)^2 &= \frac{(f \times f')^{2n}}{\gamma^{2n}}, \\ f \times f' \times \frac{(f \times f')^n}{\gamma^n} &= \frac{(f \times f')^n}{\gamma^n} \times f \times f' = \gamma \frac{(f \times f')^{n+1}}{\gamma^{n+1}} \end{aligned}$$

show that, when n becomes infinite (the convergences involved all being uni-

form), $\phi' = \phi$, $\phi \times \phi = \phi$, $f \times f' \times \phi = \phi \times f \times f' = \gamma\phi$. Finally,

$$\frac{\Gamma_n}{\gamma^n} = \frac{(f \times f')^n(1)}{\gamma^n} \rightarrow \phi(1)$$

and therefore $\phi(1) = \kappa$.

THEOREM 15. *Let $f(x)$ be as before. Then there is a (finite or infinite) sequence of real numbers $\gamma_1, \gamma_2, \dots$ and a sequence of a.p. functions $\phi_1(x), \phi_2(x), \dots$, all $\neq 0$, such that $\gamma_1 > \gamma_2 > \dots > 0$, $\phi_n' = \phi_n$, $\phi_n \times \phi_n = \phi_n$, $\phi_n(1) \geq 1$, $\phi_m \times \phi_n = 0$ ($m \neq n$), $f \times f' \times \phi_n = \phi_n \times f \times f' = \gamma_n \phi_n$, and $\gamma_1 \phi_1(x) + \gamma_2 \phi_2(x) + \dots$ converges uniformly to $f \times f'(x)$.*

Apply Theorem 14, and put there $\gamma = \gamma_1$, $\phi(x) = \phi_1(x)$. $\phi_1(1) = \kappa \geq 1$ proves that $\phi_1(x) \neq 0$. Now put $f^* = f - \phi_1 \times f$. Then $f^*(x)$ is a.p. and

$$\begin{aligned} f^* \times f^* &= (f - \phi_1 \times f) \times (f' - f' \times \phi_1) \\ &= f \times f' - \phi_1 \times f \times f' - f \times f' \times \phi_1 + \phi_1 \times f \times f' \times \phi_1 \\ &= f \times f' - \gamma_1 \phi_1 - \gamma_1 \phi_1 + \gamma_1 \phi_1 \times \phi_1 = f \times f' - \gamma_1 \phi_1. \end{aligned}$$

If $f^* \equiv 0$, this shows that $f \times f' = \gamma_1 \phi_1$, that is, the Theorem holds for a sequence consisting of one element. Assume that $f^* \neq 0$.

Then $f \times f' \times \phi_1 = \phi_1 \times f \times f' = \gamma_1 \phi_1$ implies that $(f^* \times f^{*'})^n = (f \times f' - \gamma_1 \phi_1)^n = (f \times f')^n - \gamma_1^n \phi_1$, and thus

$$\frac{(f^* \times f^{*'})^n(x)}{\gamma_1^n} = \frac{(f \times f')^n(x)}{\gamma_1^n} - \phi_1(x) \rightarrow 0 \text{ as } n \rightarrow \infty.$$

Therefore if we form $\gamma = \gamma_2$ and $\phi_2(x)$ of Theorem 14, for which we have

$$\frac{(f^* \times f^{*'})^n(x)}{\gamma_2^n} \rightarrow \phi_2(x) \neq 0,$$

then it must be the case that $\gamma_1 > \gamma_2$.

By repeating this process with $f^{**} = f^* - \phi_2 \times f^*$, $f^{***} = f^{**} - \phi_3 \times f^{**}$, $\dots$ we finally find a sequence of real numbers $\gamma_1, \gamma_2, \dots$ and two sequences of a.p. functions $\phi_1(x), \phi_2(x), \dots$ and $f(x), f^*(x), \dots$ with the following properties:

$$\gamma_1 > \gamma_2 > \dots > 0, \phi_n' = \phi_n, \phi_n \times \phi_n = \phi_n, \phi_n(1) \geq 1,$$

$$f^{(n-1)} \times f^{(n-1)'} \times \phi_n = \phi_n \times f^{(n-1)} \times f^{(n-1)'} = \gamma_n \phi_n, f^{(n)} = f^{(n-1)} - \phi_n \times f^{(n-1)},$$

these sequences ending when an $f^{(n)}$ becomes $\equiv 0$, otherwise never ending.

These rules again imply the relations $f^{(n)} \times f^{(n)'} = f^{(n-1)} \times f^{(n-1)'} - \gamma_n \phi_n$. By adding these relations for all $n = 1, \dots, p$, we obtain

$$(*) \quad \gamma_1 \phi_1 + \dots + \gamma_p \phi_p = f \times f' - f^{(p)} \times f^{(p)'}$$

We now wish to prove that $\phi_m \times \phi_n = 0$ for $m \neq n$. Application of (*) shows that it is sufficient to consider $m > n$, that is, it is sufficient to prove that $\phi_{n+k+1} \times \phi_n = 0$ for $k=0, 1, 2, \dots$. Consider the equation $f^{(n+k)} \times f^{(n+k)'} \times \phi_n = 0$. For $k=0$, the condition

$$f^{(n)} \times f^{(n)'} \times \phi_n = f^{(n-1)} \times f^{(n-1)'} \times \phi_n - \gamma_n \phi_n \times \phi_n = \gamma_n \phi_n - \gamma_n \phi_n = 0$$

obtains. If it holds for a certain $k=0, 1, 2, \dots$, we have

$$\phi_{n+k+1} \times f^{(n+k)} \times f^{(n+k)'} \times \phi_n = \begin{cases} \phi_{n+k+1} \times (f^{(n+k)} \times f^{(n+k)'} \times \phi_n) = 0, \\ (\phi_{n+k+1} \times f^{(n+k)} \times f^{(n+k)'}) \times \phi_n = \gamma_{n+k+1} \phi_{n+k+1} \times \phi_n, \end{cases}$$

and thus $\phi_{n+k+1} \times \phi_n = 0$. This gives

$$f^{(n+k+1)} \times f^{(n+k+1)'} \times \phi_n = (f^{(n+k)} \times f^{(n+k)'} - \gamma_{n+k+1} \phi_{n+k+1}) \times \phi_n = 0,$$

that is, our equation holds for $k+1$. Therefore it holds for every $k=0, 1, 2, \dots$, and with it, its consequence $\phi_{n+k+1} \times \phi_n = 0$.

Application of $\phi_n \times \dots$ or $\dots \times \phi_n$ to (*) with $p=n-1$ gives

$$\phi_n \times f \times f' = \phi_n \times f^{(n-1)} \times f^{(n-1)'} = \gamma_n \phi_n, \quad f \times f' \times \phi_n = f^{(n-1)} \times f^{(n-1)'} \times \phi_n = \gamma_n \phi_n.$$

The only thing remaining to be proved is the uniform convergence of $\gamma_1 \phi_1(x) + \dots + \gamma_n \phi_n(x)$ to $f \times f'(x)$ as $n \rightarrow \infty$, or, according to (*), the uniform convergence of $f^{(n)} \times f^{(n)'}(x)$ to 0. Now (*) implies, by (3) and (5) of Theorem 12, that $\gamma_1 \phi_1(1) + \dots + \gamma_n \phi_n(1) = (Nf)^2 - (Nf^{(n)})^2 \leq (Nf)^2$, and since $|\phi_n \times \phi_n(x)| \leq \phi_n \times \phi_n(1)$, that is, $|\phi_n(x)| \leq \phi_n(1)$, (*) implies the uniform convergence of $\gamma_1 \phi_1(x) + \gamma_2 \phi_2(x) + \dots$ to $g(x)$ as $n \rightarrow \infty$, where $g(x)$ must be a.p. Hence $f^{(n)} \times f^{(n)'}(x) \rightarrow f \times f'(x) - g(x)$ uniformly. Moreover, the above mentioned convergence implies that $\gamma_n \rightarrow 0$ (because $\phi_n(1) \geq 1$). On the other hand, we have $\Gamma_2/\Gamma_1 < \gamma$, $\Gamma_2 < \gamma \Gamma_1$, $(N[f \times f'])^2 < \gamma(Nf)^2$. If we replace f and γ by $f^{(n)}$ and γ_n we obtain $(N[f^{(n)} \times f^{(n)'}])^2 < \gamma_n(Nf^{(n)})^2 \leq \gamma_n(Nf)^2$. Thus

$$N[f^{(n)} \times f^{(n)'}] \rightarrow 0,$$

implying that $N[f \times f' - g] = 0$ and $g = f \times f'$.

9. Having reached the final result of the E. Schmidt-Weyl theory, we now pass to the approximation theorems. But as we have already mentioned, we are now giving only their proofs and shall discuss their real meaning in the next part.

DEFINITION 8. An a.p. function $\phi(x)$ such that $\phi' = \phi$ and $\phi \times \phi = \phi$ is called a unit. Two units $\phi(x)$ and $\psi(x)$ such that $\phi \times \psi = 0$ are called orthogonal.

THEOREM 16†. For every a.p. function $f(x)$ and every $\epsilon > 0$ there exists a unit $\phi(x)$ such that $N[f - \phi \times f] \leq \epsilon$.

If $f \equiv 0$, then $\phi \equiv 0$. Hence we may assume that $f \not\equiv 0$. Then apply Theorem 13. $\psi_n = \phi_1 + \cdots + \phi_n$ is a unit (because $\phi_n' = \phi_n$, $\phi_n \times \phi_n = \phi_n$, $\phi_m \times \phi_n = 0$ for $m \neq n$), and we have

$$\begin{aligned} (N[f - \psi_n \times f])^2 &= \left(f - \sum_{\nu=1}^n \phi_\nu \times f\right) \times \left(f' - \sum_{\nu=1}^n f' \times \phi_\nu\right) \quad (1) \\ &= \left(f \times f - \sum_{\nu=1}^n \phi_\nu \times f \times f - \sum_{\nu=1}^n f \times f' \times \phi_\nu + \sum_{\mu, \nu=1}^n \phi_\mu \times f \times f' \times \phi_\nu\right) \quad (1) \\ &= \left(f \times f - \sum_{\nu=1}^n \gamma_\nu \phi_\nu - \sum_{\nu=1}^n \gamma_\nu \phi_\nu + \sum_{\mu, \nu=1}^n \gamma_\nu \phi_\mu \times \phi_\nu\right) \quad (1) \\ &= \left(f \times f - \sum_{\nu=1}^n \gamma_\nu \phi_\nu\right) \quad (1) \rightarrow 0 \text{ as } n \rightarrow \infty. \end{aligned}$$

Thus $\phi = \psi_n$, for a sufficiently large n , yields the desired result.

THEOREM 17. For every a.p. function $f(x)$ and every $\epsilon > 0$ there exists an a.p. function $g(x)$ such that $|f(x) - g \times f(x)| \leq \epsilon$ for every x .

Following N. Wiener, consider the "translation function" of $f(x)$,

$$e(x) = \text{l.u.b.}_y |f(x^{-1}y) - f(y)|,$$

which was introduced by S. Bochner [2]. As $f(x)$ is a.p., it is easily seen that $e(x)$ is also a.p. Furthermore, $e(x) \geq 0$, $e(1) = 0$. Now define the function

$$F(u) = \begin{cases} 1 - \frac{u}{\epsilon}, & 0 \leq u \leq \epsilon, \\ 0, & u \geq \epsilon. \end{cases}$$

As $F(u)$ is continuous, $\phi(x) = F(e(x))$ is a.p. It is obvious that $\phi(x) \geq 0$, $\phi(1) = 1$; and if $\text{l.u.b.}_y |f(x^{-1}y) - f(y)| > \epsilon$, then $\phi(x) = 0$. Thus $\phi(x) \neq 0$ implies that $|f(x^{-1}y) - f(y)| \leq \epsilon$, and therefore, always $|\phi(x)(f(x^{-1}y) - f(y))| \leq \epsilon \phi(x)$. Consequently, $|M_x[\phi(x)(f(x^{-1}y) - f(y))]| \leq \epsilon M_x \phi(x)$. Now $M_x[\phi(x)(f(x^{-1}y) - f(y))] = \phi \times f(y) - f(y)M_x \phi(x)$, and therefore the function

† Cf. Theorems 28 and 29, where the results of Theorems 16 and 18 will be interpreted and applied.

$$g(y) = \frac{\phi(y)}{M_x \phi(x)}$$

meets the requirements.

THEOREM 18. *For every a.p. function $f(x)$ and every $\epsilon > 0$ there exists an a.p. function $g(x)$ and a unit $\phi(x)$ such that $|f(x) - \phi \times g \times f(x)| \leq \epsilon$ for every x .†*

Choose the function $g(x)$ of Theorem 17 corresponding to $f(x)$ and $\epsilon/2$, and choose the function $\phi(x)$ of Theorem 16 corresponding to $g(x)$ and $\epsilon/(2Nf)$. Then $|f(x) - g \times f(x)| \leq \epsilon/2$, $|g \times f(x) - \phi \times g \times f(x)| \leq N[g - \phi \times g]Nf \leq \epsilon/2$, and therefore $|f(x) - \phi \times g \times f(x)| \leq \epsilon$.

III. THEORY OF LINEAR REPRESENTATIONS OF $\mathfrak{G}$

10. We define the representations in the usual manner:

DEFINITION 9. *If to every $a \in \mathfrak{G}$ there corresponds a matrix*

$$D(a) = \{D_{\rho\sigma}(a)\} \quad (\rho, \sigma = 1, \dots, s)$$

of degree s such that $D(1) = 1$, (the unit matrix of degree s), $D(ab) = D(a) \cdot D(b)$, then we call $D(a)$ a representation of $\mathfrak{G}$. (No continuity is assumed.) Two representations $D(a)$ and $D'(a)$ are called equivalent if they are of the same degree s , and if a fixed matrix $U = \{U_{\rho\sigma}\}$, $\rho, \sigma = 1, \dots, s$, exists which transforms one representation into the other: $U^{-1}D(a)U = D'(a)$. A representation $D(a)$ is called reducible (completely reducible) if it is equivalent to a representation $D'(a)$ such that $D'_{\rho\sigma}(a) = 0$ identically (in a) whenever $\rho \leq t$, $\sigma > t$ ($\rho \leq t$, $\sigma > t$ or $\rho > t$, $\sigma \leq t$)‡, for a fixed value of t , $1 \leq t \leq s-1$. Representations without these properties are irreducible (completely irreducible).

For finite groups $\mathfrak{G}$ Frobenius and Schur gave a complete theory of all representations [21, 22]; for continuous groups $\mathfrak{G}$ close analogues of their results were established by Schur for the rotation group in three dimensions, and in much broader generality by Weyl for all compact Lie-groups [30]. These results were extended to all compact groups $\mathfrak{G}$ by Haar [11, pp. 166–169] with the help of his notion of “right-invariant” Lebesgue measure in groups. We shall push the extension further to all groups $\mathfrak{G}$, but in order to do this it is natural and necessary to restrict the domain of representations of $\mathfrak{G}$ by means of

† Cf. footnote to Theorem 16.

‡ These are the fundamental notions of the Frobenius-Schur theory of group representations [21, 22, 30].

THEOREM 19. *The following conditions on a representation $D(a)$ of $\mathfrak{G}$ are equivalent to each other:*

A. $D(a)$ is equivalent to a unitary† representation.

B. All elements $D_{\rho\sigma}(a)$ of $D(a)$ are bounded.

C. All elements $D_{\rho\sigma}(a)$ of $D(a)$ are a.p.

If $D'(a)$ is unitary, then, by the footnote just cited in the case where $\rho = \sigma$,

$$\sum_{\tau=1}^s |D_{\rho\tau}'(a)|^2 = 1, \quad |D_{\rho\tau}'(a)| \leq 1,$$

and all $D_{\rho\sigma}'(a)$ are bounded. Therefore the elements $D_{\rho\tau}(a)$ of any $D(a)$ equivalent to $D'(a)$ must also be bounded. Thus A implies B.

If all $D_{\rho\sigma}(a)$ are bounded, every sequence $D_{\rho\sigma}(a_n)$, $n=1, 2, \dots$, contains a subsequence which converges for all $\rho, \sigma=1, \dots, s$. And then, since $D(xa_n) = D(x)D(a_n)$ and $D(a_nx) = D(a_n)D(x)$, the representations $D(xa_n)$ and $D(a_nx)$, and hence all $D_{\rho\sigma}(xa_n)$ and $D_{\rho\sigma}(a_nx)$, converge uniformly. Thus all $D_{\rho\sigma}(a)$ are a.p., that is, B implies C.

If all $D_{\rho\sigma}(a)$ are a.p., so are the expressions $\sum_{\tau=1}^s D_{\rho\tau}(a) \overline{D_{\sigma\tau}(a)}$, and we can form

$$A_{\rho\sigma} = M_x \left[\sum_{\tau=1}^s D_{\rho\tau}(x) \overline{D_{\sigma\tau}(x)} \right].$$

Now

$$\sum_{\tau=1}^s D_{\rho\tau}(a) \overline{D_{\sigma\tau}(a)} = \overline{\sum_{\tau=1}^s D_{\sigma\tau}(a) \overline{D_{\rho\tau}(a)}}$$

and for every system $\xi_1, \dots, \xi_s$ which is not identically zero,

$$\sum_{\rho, \sigma=1}^s \left[\sum_{\tau=1}^s D_{\rho\tau}(a) \overline{D_{\sigma\tau}(a)} \right] \xi_\rho \bar{\xi}_\sigma = \sum_{\tau=1}^s \left| \sum_{\rho=1}^s D_{\rho\tau}(a) \xi_\rho \right|^2 > 0,$$

so that

$$A_{\rho\sigma} = \overline{A_{\sigma\rho}} \quad \text{and} \quad \sum_{\rho, \sigma=1}^s A_{\rho\sigma} \xi_\rho \bar{\xi}_\sigma > 0.$$

Therefore the matrix $A = \{A_{\rho\sigma}\}$ is Hermitian and positive definite. Hence

† That is, to a representation $D'(a)$ in which all matrices $D'(a) = \{D'_{\rho\sigma}(a)\}$ ($\rho, \sigma=1, \dots, s$) are unitary. A matrix $U = \{U_{\rho\sigma}\}$ is called unitary if its adjoint $U^* = \{\bar{U}_{\sigma\rho}\}$ is reciprocal to it, that is, if $UU^* = U^*U = 1$, or more explicitly,

$$\sum_{\tau=1}^s U_{\rho\tau} \bar{U}_{\sigma\tau} = \sum_{\tau=1}^s U_{\tau\rho} \bar{U}_{\tau\sigma} = \delta_{\rho\sigma} = \begin{cases} 1, & \rho = \sigma, \\ 0, & \rho \neq \sigma. \end{cases}$$

there exists a matrix $X = \{X_{\rho\sigma}\}$ such that $A = XX^*$. On the other hand,

$$\begin{aligned} A_{\rho\sigma} &= M_x \left[\sum_{\tau=1}^i D_{\rho\tau}(x) \overline{D_{\sigma\tau}(x)} \right] = M_x \left[\sum_{\tau=1}^i D_{\rho\tau}(ax) \overline{D_{\sigma\tau}(ax)} \right] \\ &= M_x \left[\sum_{\rho', \sigma'=1}^i \sum_{\tau=1}^i D_{\rho\rho'}(a) D_{\rho'\tau}(x) \overline{D_{\sigma\sigma'}(a)} \overline{D_{\sigma'\tau}(x)} \right] \\ &= \sum_{\rho', \sigma'=1}^i D_{\rho\rho'}(a) \overline{D_{\sigma\sigma'}(a)} M_x \left[\sum_{\tau=1}^i D_{\rho'\tau}(x) \overline{D_{\sigma'\tau}(x)} \right] \\ &= \sum_{\rho', \sigma'=1}^i D_{\rho\rho'}(a) \overline{D_{\sigma\sigma'}(a)} A_{\rho'\sigma'}, \end{aligned}$$

that is, $A = D(a)AD(a)^*$, or $XX^* = D(a)XX^*D(a)^*$, $X^{-1}D(a)XX^*D(a)^*X^{*-1} = 1$, $(X^{-1}D(a)X)(X^{-1}D(a)X)^* = 1$. In other words, the equivalent representation $X^{-1}D(a)X = D'(a)$ is unitary. Thus C implies A.

Our three statements together prove the equivalence of A, B, and C.

DEFINITION 10. We call *normal* the representations satisfying one of the equivalent conditions of Theorem 19.

11. The fundamental theorems of the theory of orthogonal representations may now be proved in the classical way [21, 22, 30, 32].

THEOREM 20. Let $D(a)$ and $E(a)$ be completely irreducible normal representations of degrees s and t respectively, and let A be a rectangular matrix with s rows and t columns. If $D(a)A \equiv AE(a)$ for every a , then either $A = 0$ or $s = t$ and $\det A \neq 0$, the latter alternative of course implying the equivalence of $D(a)$ and $E(a)$. If $D(a) = E(a)$, then $A = \alpha 1$ (α being a complex number).

In all these statements (except the last) $D(a)$ and $E(a)$ may be replaced by two equivalent representations. Therefore we may assume them to be unitary. Even then further transformations by unitary matrices X and Y are possible. They carry A into $A' = X^{-1}AY$. Now by such transformations we can obtain $A' = \{A'_{\rho\sigma}\}$, $\rho = 1, \dots, s$; $\sigma = 1, \dots, t$, such that

$$A'_{\rho\sigma} = \begin{cases} c_\rho, & \text{for } \rho = \sigma = 1, \dots, r, \text{ all } c_\rho > 0, \\ 0, & \text{for all other } \rho \text{ and } \sigma, \end{cases}$$

where r is the rank of A' and $r \leq s$, $r \leq t$. Therefore we may assume that A itself has this form.

Under these conditions the relation $D(a)A = AE(a)$ implies that $D_{\rho\sigma}(a) = 0$ for $\rho > r$ and $\sigma \leq r$, and that $E_{\rho\sigma}(a) = 0$ for $\rho \leq r$ and $\sigma > r$. Since A^* also has the form we assumed for A , $D(a)^* = D(a)^{-1} = D(a^{-1})$, $E(a)^* = E(a)^{-1} = E(a^{-1})$,

we get, by applying $*$ and replacing a by a^{-1} , $A^*D(a) = E(a)A^*$, so that $D_{\rho\sigma}(a) = 0$ for $\rho \leq r$ and $\sigma > r$, and $E_{\rho\sigma}(a) = 0$ for $\rho > r$ and $\sigma \leq r$. Thus the complete irreducibility of $D(a)$ requires that r be 0 or s , and the complete irreducibility of $E(a)$ requires that r be 0 or t . Hence either $r=0$, in which case $A=0$, or $r=s=t$, in which case $\det A \neq 0$.

If $D(a) = E(a)$, every $A - \alpha 1$ has the same property as A . If α is a root of the characteristic equation of A , we have $\det [A - \alpha 1] = 0$, so that our alternative requires that $A - \alpha 1 = 0$, and $A = \alpha 1$.

THEOREM 21. *Let $D(a)$ and $E(a)$ be completely irreducible normal representations of degrees s and t respectively. If they are inequivalent, we have*

$$D_{\rho\sigma} \times E_{\tau\nu}(x) \equiv 0.$$

Considering $D(a)$ alone, we have

$$D_{\rho\sigma} \times D_{\tau\nu}(x) \equiv \begin{cases} \frac{1}{s} D_{\rho\nu}(x) & \text{for } \sigma = \tau, \\ 0 & \text{for } \sigma \neq \tau. \end{cases}$$

Form the (rectangular) matrix

$$A = \{A_{\tau\sigma}\}, \quad A_{\tau\sigma} = D_{\rho\sigma} \times E_{\tau\nu}(x)$$

for a given choice of ρ, ν , and x . Then

$$\begin{aligned} \sum_{\tau'=1}^t E_{\tau\tau'}(a) A_{\tau'\sigma} &= M_y \left[\sum_{\tau'=1}^t D_{\rho\sigma}(xy^{-1}) E_{\tau\tau'}(a) E_{\tau'\nu}(y) \right] \\ &= M_y [D_{\rho\sigma}(xy^{-1}) E_{\tau\nu}(ay)] = M_z [D_{\rho\sigma}(z) E_{\tau\nu}(az^{-1}x)], \\ \sum_{\sigma'=1}^s A_{\tau\sigma'} D_{\sigma'\sigma}(a) &= M_y \left[\sum_{\sigma'=1}^s D_{\rho\sigma'}(y) D_{\sigma'\sigma}(a) E_{\tau\nu}(y^{-1}x) \right] \\ &= M_y [D_{\rho\sigma}(ya) E_{\tau\nu}(y^{-1}x)] = M_z [D_{\rho\sigma}(z) E_{\tau\nu}(az^{-1}x)] \end{aligned}$$

(the variable y being replaced by $z = xy^{-1}$ and $z = ya$ respectively), and therefore

$$E(a)A = AD(a).$$

Thus we can apply Theorem 20. If $D(a)$ and $E(a)$ are inequivalent, it results that $A=0$, and if $D(a) = E(a)$, $A = \alpha_{\rho\nu}(x) 1$, ($\alpha_{\rho\nu}(x)$ being a complex number). This implies (1) that $A_{\tau\sigma} = 0$, that is, $D_{\rho\sigma} \times E_{\tau\nu}(x) = 0$ if $D(a)$ and $E(a)$ are inequivalent or if $D(a) = E(a)$ and $\sigma \neq \tau$; and (2) that $A_{\tau\sigma} = \alpha_{\rho\nu}(x)$, which is independent of σ , if $D(a) = E(a)$ and $\sigma = \tau$. Hence all the statements of our Theorem are proved if we show that $\alpha_{\rho\nu}(x) = (1/s) D_{\rho\nu}(x)$. This follows immediately from

$$\begin{aligned} s\alpha_{\rho\nu}(x) &= \sum_{\sigma=1}^s A_{\sigma\sigma} = M_y \left[\sum_{\sigma=1}^s D_{\rho\sigma}(xy^{-1})D_{\sigma\nu}(y) \right] \\ &= M_y D_{\rho\nu}(xy^{-1}y) = M_y D_{\rho\nu}(x) = D_{\rho\nu}(x). \end{aligned}$$

THEOREM 22. *For normal representations reducibility and complete reducibility are equivalent, so that irreducibility and complete irreducibility are also equivalent.*

That complete reducibility implies reducibility is obvious. Assume now that $D(a)$ is reducible without being completely reducible. As we can replace $D(a)$ by any equivalent representation, we may assume $D(a)$ to be in the form described in Definition 9. Then there would be a pair of indices ρ and σ such that $D_{\rho\sigma}(x) \equiv 0$. By Theorem 21, this relation implies that $D_{\rho\sigma} \times D_{\sigma\rho} \equiv D_{\rho\rho} \equiv 0$, in spite of the fact that $D_{\rho\rho}(1) = 1$. Thus $D(a)$ must be completely reducible.

12. Our next task is to formulate the connection between the units of Definition 8 and representations. This is accomplished by means of

THEOREM 23. *For every unit $\phi(x)$ there exist a number of inequivalent irreducible unitary representations $D^{(1)}(a), \dots, D^{(u)}(a)$ of degrees $s_1, \dots, s_u$ respectively such that*

$$\phi(x) = \sum_{\omega=1}^u s_{\omega} \sum_{\rho, \sigma=1}^{s_{\omega}} \alpha_{\rho\sigma}^{(\omega)} D_{\rho\sigma}^{(\omega)}(x).$$

Here every matrix $\alpha^{(\omega)} = \{\alpha_{\rho\sigma}^{(\omega)}\}$ is idempotent, that is, $\alpha^{(\omega)*} = \alpha^{(\omega)}$ (cf. footnote on page 465), $(\alpha^{(\omega)})^2 = \alpha^{(\omega)}$.† Conversely, every $\phi(x)$ which is formed in this way (where $D^{(\omega)}(a)$ and $\alpha^{(\omega)}$ satisfy our conditions) is a unit.

By a suitable choice of $D^{(\omega)}(a)$ we can give the matrices $\alpha^{(\omega)}$ the form

$$\alpha_{\rho\sigma}^{(\omega)} = \begin{cases} 1 & \text{for } \rho = \sigma \leq s'_{\omega}, \text{ with some } s'_{\omega} = 1, \dots, s_{\omega}, \\ 0 & \text{for all other } \rho \text{ and } \sigma. \end{cases}$$

Consider the (a.p.) solutions $f(x)$ of the equation $\phi \times f = f$. Assume that it is possible to find s solutions $g_1, \dots, g_s$ among them which satisfy the conditions

$$g_{\mu} \times g'_{\nu}(1) = M_x [g_{\mu}(x) \overline{g'_{\nu}(x)}] = \begin{cases} 1 & \text{for } \mu = \nu, \\ 0 & \text{for } \mu \neq \nu. \end{cases}$$

Put

† That is,

$$\alpha_{\rho\sigma}^{(\omega)} = \overline{\alpha_{\sigma\rho}^{(\omega)}}, \quad \sum_{\sigma=1}^{s_{\omega}} \alpha_{\rho\sigma}^{(\omega)} \alpha_{\sigma\tau}^{(\omega)} = \alpha_{\rho\tau}^{(\omega)}.$$

$$\psi(x, y) = \phi(xy^{-1}) - \sum_{\mu=1}^s g_{\mu}(x) \overline{g_{\mu}(y)}.$$

Then

$$\begin{aligned} M_{x,y} [|\psi(x, y)|^2] &= M_{x,y} [|\phi(xy^{-1})|^2] - \sum_{\mu=1}^s M_{x,y} [\phi(xy^{-1}) \overline{g_{\mu}(x)} g_{\mu}(y)] \\ &\quad - \sum_{\mu=1}^s M_{x,y} [\overline{\phi(xy^{-1})} g_{\mu}(x) \overline{g_{\mu}(y)}] + \sum_{\mu, \nu=1}^s M_{x,y} [g_{\mu}(x) \overline{g_{\mu}(y)} \overline{g_{\nu}(x)} g_{\nu}(y)]. \end{aligned}$$

By Theorems 9 and 10, the orthogonality properties of g_{μ} , and the relations $\phi \times g_{\mu} = g_{\mu}$, $g'_{\mu} \times \phi = g'_{\mu}$, this turns out to be equal to

$$M_x [|\phi(x)|^2] - s - s + s = (N\phi)^2 - s.$$

Therefore we have $(N\phi)^2 - s \geq 0$, so that $s \leq (N\phi)^2$. Thus the possible numbers s are bounded, and they have a maximal value. Assume that s is this maximal value and choose $g_1, \dots, g_s$ accordingly. If a solution $f(x)$ of $\phi \times f = f$ is such that $f \times g'_{\mu} = 0$ for all $\mu = 1, \dots, s$, then necessarily $f \equiv 0$, for otherwise we could put $g_{s+1} = f/(Nf)$, implying that $Ng_{s+1} = 1$, that is, $g_{s+1} \times g'_{s+1}(1) = 1$, and $g_{s+1} \times g'_{\mu}(1) = 0$ as well as $g_{\mu} \times g'_{s+1}(1) = 0$ for $\mu = 1, \dots, s$, so that

$$g_{\mu} \times g'_{\nu}(1) = \begin{cases} 1 & \text{for } \mu = \nu, \\ 0 & \text{for } \mu \neq \nu, \end{cases} \quad \mu, \nu = 1, \dots, s+1,$$

contradicting our assumption that s is maximal.

If we define $f(x)$ to be $\overline{\psi(x, a)}$, we find by a simple computation that $\phi \times f = f$ and $f \times g'_{\mu}(1) = 0$. Therefore $f(x) \equiv 0$, $\psi(x, a) \equiv 0$, and, as a was arbitrary, $\psi(x, y) \equiv 0$, that is,

$$(\dagger) \quad \phi(xy^{-1}) \equiv \sum_{\mu=1}^s g_{\mu}(x) \overline{g_{\mu}(y)}.$$

As $(\dagger)$ implies that

$$\phi \times f(x) = M_y [\phi(xy^{-1})f(y)] = \sum_{\mu=1}^s M_y [\overline{g_{\mu}(y)}f(y)]g_{\mu}(x),$$

every solution of $\phi \times f = f$ has the form $\sum_{\mu=1}^s \alpha_{\mu} g_{\mu}$ (α_{μ} being complex numbers). It is obvious that $f(xa)$ is a solution along with $f(x)$, so that $g_{\mu}(xa)$ is a solution, and we can write

$$g_{\sigma}(xa) = \sum_{\rho=1}^s D_{\rho\sigma}(a)g_{\rho}(x).$$

The orthogonality properties of g_μ determine the coefficients $D_{\rho\sigma}(a)$ uniquely:

$$D_{\rho\sigma}(a) = M_s[g_\sigma(xa)\overline{g_\rho(x)}] = g'_\rho \times g_\sigma(a).$$

Hence $D(1) = 1_s$, and if we put $D(a) = \{D_{\rho\sigma}(a)\}$, $\rho, \sigma = 1, \dots, s$, we obviously have $D(ab) = D(a)D(b)$, that is, $D(a)$ is a representation. As (§) holds if we replace x and y by xa and ya , that is, if we replace $g_\mu(x)$ and $g_\mu(y)$ by $g_\mu(xa)$ and $g_\mu(ya)$, the transformations $D(a)$ must be unitary.

All this implies that

$$\phi(x) = \sum_{\sigma=1}^s g_\sigma(x)\overline{g_\sigma(1)} = \sum_{\rho, \sigma=1}^s D_{\rho\sigma}(x)g_\rho(1)\overline{g_\sigma(1)},$$

so that $\phi(x)$ is a linear aggregate of all $D_{\rho\sigma}(x)$. Now (cf. Theorem 22) $D(a)$ can be transformed into an equivalent $D'(a)$ which consists of a certain number of irreducible representations $D^{(1)}(a), \dots, D^{(v)}(a)$ (of degrees $s_1, \dots, s_v$ respectively, where $s_1 + \dots + s_v = s$) which succeed each other along the main diagonal, and zeros in all other places. Therefore $\phi(x)$ is also a linear aggregate of the elements $D_{\rho\sigma}'(x)$, that is, of the elements $D_{\rho\sigma}^{(\omega)}(x)$. Now if some representation $D^{(\omega)}(x)$ is equivalent to another representation $D^{(\pi)}(x)$, the elements $D_{\rho\sigma}^{(\omega)}(x)$ are linear aggregates of elements of $D^{(\pi)}(x)$. Therefore, in expressing $\phi(x)$ as a linear aggregate of all elements $D_{\rho\sigma}^{(\omega)}(x)$, it is sufficient to keep only one member of each class of equivalent representations $D^{(\omega)}(x)$. Those representations thus kept may be labeled $D^{(1)}(x), \dots, D^{(u)}(x)$, $u \leq v$. So we finally have the result that $D^{(1)}(a), \dots, D^{(u)}(a)$ (of degrees $s_1, \dots, s_u$ respectively) is a set of inequivalent irreducible unitary representations, and $\phi(x)$ is a linear aggregate of the elements $D_{\rho\sigma}^{(\omega)}(x)$, that is, we can write $\phi(x)$ in the form

$$\phi(x) = \sum_{\omega=1}^u s_\omega \sum_{\rho, \sigma=1}^{s_\omega} \alpha_{\rho\sigma}^{(\omega)} D_{\rho\sigma}^{(\omega)}(x).$$

If by means of this equation we now determine the meaning of $\phi' = \phi$ and $\phi \times \phi = \phi$, remembering that

$$(D^{(\omega)}(x))^* = D^{(\omega)}(x)^{-1} = D^{(\omega)}(x^{-1}),$$

$$\overline{D_{\sigma\rho}^{(\omega)}(x)} = D_{\rho\sigma}^{(\omega)}(x^{-1}), \quad \overline{D_{\rho\sigma}^{(\omega)}(x^{-1})} = D_{\sigma\rho}^{(\omega)}(x),$$

that is, $D_{\sigma\rho}^{(\omega)'}(x) = D_{\sigma\rho}^{(\omega)}(x)$, and

$$D_{\rho\sigma}^{(\omega)} \times D_{\tau\nu}^{(\chi)} = \begin{cases} \frac{1}{s_\omega} D_{\rho\nu}^{(\omega)} & \text{if } \omega = \chi \text{ and } \sigma = \tau, \\ 0 & \text{for all other } \omega, \chi, \rho, \sigma, \tau, \nu, \end{cases}$$

we obtain exactly the conditions in our Theorem. Furthermore it is clear that every matrix $\alpha^{(\omega)} = \{\alpha_{\rho\sigma}^{(\omega)}\}$, being idempotent, can be transformed into the form given at the end of our Theorem. And the inverse transformations of the representations $D^{(\omega)}(x)$, which carry them into equivalent representations, bring about just these $\alpha^{(\omega)}$ -transformations.

13. We choose a system of "representants" for the inequivalent irreducible (normal or orthogonal) representations of $\mathfrak{G}$:

DEFINITION 11. Let I be the set of all irreducible normal representations of $\mathfrak{G}$. Call each subset $\mathfrak{E}$ of I which consists of all the elements of I equivalent to one of its elements a class. It is obvious that every element of I belongs to exactly one class. Call the set of all classes C . Each element $\mathfrak{E}$ of C contains unitary representations (since every normal representation is equivalent to a unitary representation). Select one unitary representation from each $\mathfrak{E}$ of C , call it the representant of $\mathfrak{E}$, denote it by $D(a; \mathfrak{E})^\dagger$, and denote its degree by $s(\mathfrak{E})$.

THEOREM 24. The (a.p.) functions $D_{\rho\sigma}(x; \mathfrak{E})$, $\mathfrak{E}$ in C , ρ and $\sigma = 1, \dots, s(\mathfrak{E})$, have the property that

$$D_{\rho\sigma}(\mathfrak{E}) \times D_{\tau\nu}(\mathfrak{D}) = \begin{cases} \frac{1}{s(\mathfrak{E})} D_{\rho\nu}(\mathfrak{E}) & \text{for } \mathfrak{E} = \mathfrak{D} \text{ and } \sigma = \tau, \\ 0 & \text{for all other } \mathfrak{E}, \mathfrak{D}, \rho, \sigma, \tau, \nu. \end{cases}$$

This implies that

$$M_x[D_{\rho\sigma}(x; \mathfrak{E}) \overline{D_{\tau\nu}(x; \mathfrak{D})}] = \begin{cases} \frac{1}{s(\mathfrak{E})} & \text{for } \mathfrak{E} = \mathfrak{D}, \rho = \nu, \sigma = \tau, \\ 0 & \text{for all other } \mathfrak{E}, \mathfrak{D}, \rho, \sigma, \tau, \nu. \end{cases}$$

The first formula has been proved in Theorem 21. The second formula follows from it if we put the variable equal to 1 and remember that $D_{\tau\tau}(\mathfrak{D})' = \overline{D_{\tau\tau}(\mathfrak{D})}$, $D(1, \mathfrak{D}) = 1$.

Thus the functions $s(\mathfrak{E})^{1/2} D_{\rho\sigma}(x; \mathfrak{E})$ form a "normalized orthogonal" system. This is the basis for the formulation of the usual expansion theorems. The key theorems of the theory will be proved as Theorems 28 and 29.

[†] This does not imply an essential use of the "axiom of choice" because it would in most cases be possible to characterize a $D(a; \mathfrak{E})$ in $\mathfrak{E}$ in a unique way. To abbreviate we shall also use the notation $D(\mathfrak{E})$ omitting the argument a .

DEFINITION 12. If $f(x)$ is an a.p. function, the complex numbers

$$\tilde{\alpha}_{\rho\sigma}(\mathbb{E}) = f \times D_{\rho\sigma}(1; \mathbb{E})' = M_z[f(x)\overline{D_{\rho\sigma}(x; \mathbb{E})}]$$

are called its expansion coefficients. The matrices $\tilde{\alpha}(\mathbb{E}) = \{\tilde{\alpha}_{\rho\sigma}(\mathbb{E})\}$ are called the expansion matrices.

THEOREM 25. If f and g have the expansion matrices $\tilde{\alpha}(\mathbb{E})$ and $\tilde{\beta}(\mathbb{E})$ ($\mathbb{E}$ running over C), $f \pm g$, θf (θ any complex number), f' and $f \times g$ have the expansion matrices $\tilde{\alpha}(\mathbb{E}) \pm \tilde{\beta}(\mathbb{E})$, $\theta \tilde{\alpha}(\mathbb{E})$, $\tilde{\alpha}(\mathbb{E})^*$, and $\tilde{\alpha}(\mathbb{E})\tilde{\beta}(\mathbb{E})$ respectively. A unit ϕ is characterized by the fact that all its expansion matrices are idempotent, so that only a finite number of them can be $\neq 0$.

The statements concerning $f \pm g$ and θf are obvious; as to f' it is sufficient to remark that

$$\begin{aligned} M_z[f'(x)\overline{D_{\rho\sigma}(x; \mathbb{E})}] &= M_z[\overline{f(x^{-1})D_{\rho\sigma}(x^{-1}; \mathbb{E})}'] \\ &= M_z[\overline{f(x)D_{\rho\sigma}(x; \mathbb{E})}'] = \overline{M_z[f(x)D_{\rho\sigma}(x; \mathbb{E})]} = \overline{\tilde{\alpha}_{\rho\sigma}(\mathbb{E})}. \end{aligned}$$

The following computation proves our statement with regard to $f \times g$ (cf. Theorems 7 and 8):

$$\begin{aligned} M_z[f \times g(x)\overline{D_{\rho\sigma}(x; \mathbb{E})}] &= M_{z,y}[f(y)g(y^{-1}x)\overline{D_{\rho\sigma}(x; \mathbb{E})}] \\ &= M_{y,z}[f(y)g(z)\overline{D_{\rho\sigma}(yz; \mathbb{E})}] \\ &= \sum_{\tau=1}^{s(\mathbb{E})} M_{y,z}[f(y)g(z)\overline{D_{\rho\tau}(y; \mathbb{E})} \overline{D_{\tau\sigma}(z; \mathbb{E})}] \\ &= \sum_{\tau=1}^{s(\mathbb{E})} M_y[f(y)\overline{D_{\rho\tau}(y; \mathbb{E})}] M_z[g(z)\overline{D_{\tau\sigma}(z; \mathbb{E})}] \\ &= \sum_{\tau=1}^{s(\mathbb{E})} \tilde{\alpha}_{\rho\tau}(\mathbb{E})\tilde{\beta}_{\tau\sigma}(\mathbb{E}). \end{aligned}$$

This discussion shows that the idempotence of all expansion matrices $\tilde{\alpha}(\mathbb{E})$ is characteristic of units ϕ , but it follows from Theorem 21 that all those matrices which have a $D(a; \mathbb{E})$ not identical to a $D^{(\omega)}(a)$ for some $\omega = 1, \dots, u$ must vanish, so that only a finite number are different from zero. (We here use the fact that if two functions have the same expansion matrices, they coincide; that is, if all expansion matrices of a function f vanish, then $f(x) \equiv 0$). This follows from Theorem 28 (the proof of which does not depend upon Theorem 25) by putting $f(x) \equiv g(x)$.

THEOREM 26. If f has the expansion matrices $\tilde{\alpha}(\mathbb{E}) = \{\tilde{\alpha}_{\rho\sigma}(\mathbb{E})\}$ and if $\mathbb{E}_1, \dots, \mathbb{E}_n$ are elements of C , then of all linear aggregates g of the elements $D_{\rho\sigma}(\mathbb{E}_\omega)$, $\omega = 1, \dots, n$; $\rho, \sigma = 1, \dots, s(\mathbb{E}_\omega)$, which can be written in the form

$$g(x) = \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} \alpha_{\rho\sigma}(\mathbb{G}_{\omega}) D_{\rho\sigma}(x; \mathbb{G}_{\omega})$$

that one which minimizes the expression $(N[f-g])^2$ is characterized by the property that $\alpha_{\rho\sigma}(\mathbb{G}_{\omega}) = \tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega})$. The minimum value in question is

$$(Nf)^2 - \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega})|^2.$$

The proof is contained in the well known computation

$$\begin{aligned} (N[f-g])^2 &= M_x[|f(x) - g(x)|^2] \\ &= M_x[|f(x)|^2] - M_x[f(x)\overline{g(x)}] - M_x[\overline{f(x)}g(x)] + M_x[|g(x)|^2] \\ &= (Nf)^2 - \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} \overline{\alpha_{\rho\sigma}(\mathbb{G}_{\omega})} M_x[f(x)\overline{D_{\rho\sigma}(x; \mathbb{G}_{\omega})}] \\ &\quad - \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} \alpha_{\rho\sigma}(\mathbb{G}_{\omega}) M_x[\overline{f(x)} D_{\rho\sigma}(x; \mathbb{G}_{\omega})] \\ &\quad + \sum_{\omega, \chi=1}^n s(\mathbb{G}_{\omega}) s(\mathbb{G}_{\chi}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} \sum_{\tau, \nu=1}^{s(\mathbb{G}_{\chi})} \alpha_{\rho\sigma}(\mathbb{G}_{\omega}) \overline{\alpha_{\tau\nu}(\mathbb{G}_{\chi})} M_x[D_{\rho\sigma}(x; \mathbb{G}_{\omega}) \overline{D_{\tau\nu}(x; \mathbb{G}_{\chi})}] \\ &= (Nf)^2 - \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} \tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega}) \overline{\alpha_{\rho\sigma}(\mathbb{G}_{\omega})} \\ &\quad - \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} \overline{\tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega})} \alpha_{\rho\sigma}(\mathbb{G}_{\omega}) + \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} |\alpha_{\rho\sigma}(\mathbb{G}_{\omega})|^2 \\ &= \left\{ (Nf)^2 - \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega})|^2 \right\} \\ &\quad + \sum_{\omega=1}^n s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} |\alpha_{\rho\sigma}(\mathbb{G}_{\omega}) - \tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega})|^2. \end{aligned}$$

THEOREM 27. (Bessel's inequality.) If f has the expansion matrices $(\tilde{\alpha}\mathbb{G}) = \{\tilde{\alpha}_{\rho\sigma}(\mathbb{G})\}$, the number of those $\mathbb{G}$ for which $\tilde{\alpha}(\mathbb{G}) \neq 0$ is at most countably infinite, that is, it is possible to arrange them in a finite or infinite sequence $\mathbb{G}_1, \mathbb{G}_2, \dots$. Then we have

$$(Nf)^2 \geq \sum_{\omega} s(\mathbb{G}_{\omega}) \sum_{\rho, \sigma=1}^{s(\mathbb{G}_{\omega})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{G}_{\omega})|^2.$$

Since, for all other $\mathbb{G}$'s, $\tilde{\alpha}_{\rho\sigma}(\mathbb{G}) = 0$, we can instead write

$$(Nf)^2 \geq \sum_{\mathbb{G}} s(\mathbb{G}) \sum_{\rho, \sigma=1}^{s(\mathbb{G})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{G})|^2.$$

The number of $\mathbb{C}$'s for which

$$s(\mathbb{C}) \sum_{\rho, \sigma=1}^{s(\mathbb{C})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{C})|^2 \geq \epsilon, \quad \epsilon > 0,$$

is, because of the last statement of Theorem 26, certainly finite and $\leq (Nf)^2/\epsilon$. Putting $\epsilon=1, \frac{1}{2}, \frac{1}{3}, \dots$ successively, we see that, with at most countably infinitely many exceptional $\mathbb{C}$'s, it is always the case that

$$s(\mathbb{C}) \sum_{\rho, \sigma=1}^{s(\mathbb{C})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{C})|^2 = 0,$$

that is, $\tilde{\alpha}_{\rho\sigma}(\mathbb{C})=0$. This proves the first statement of our Theorem.

For every n the last statement of Theorem 26 shows that

$$\sum_{\omega=1}^n s(\mathbb{C}_\omega) \sum_{\rho, \sigma=1}^{s(\mathbb{C}_\omega)} |\tilde{\alpha}_{\rho\sigma}(\mathbb{C}_\omega)|^2 \leq (Nf)^2.$$

Hence the sum

$$\sum_{\omega} s(\mathbb{C}_\omega) \sum_{\rho, \sigma=1}^{s(\mathbb{C}_\omega)} |\tilde{\alpha}_{\rho\sigma}(\mathbb{C}_\omega)|^2$$

is convergent (if the number of terms is infinite) and $\leq (Nf)^2$.

THEOREM 28. (*Parseval's equation.*) *If f and g have the expansion matrices $\tilde{\alpha}(\mathbb{C}) = \{\tilde{\alpha}_{\rho\sigma}(\mathbb{C})\}$ and $\tilde{\beta}(\mathbb{C}) = \{\tilde{\beta}_{\rho\sigma}(\mathbb{C})\}$, then*

$$M_z[f(x)\overline{g(x)}] = \sum_{\mathbb{C}} s(\mathbb{C}) \sum_{\rho, \sigma=1}^{s(\mathbb{C})} \tilde{\alpha}_{\rho\sigma}(\mathbb{C}) \overline{\tilde{\beta}_{\rho\sigma}(\mathbb{C})},$$

where the series $\sum_{\mathbb{C}}$ contains at most countably infinitely many terms $\neq 0$, and is absolutely convergent (if infinite at all).

If this is proved for $f=g$, we obtain the real part of the statement by replacing our f by $(f+g)/2$ and $(f-g)/2$ and subtracting. Replacing f and g by if and g gives the imaginary part and completes the proof. Hence we may assume that $f=g$, that is, we must show that

$$(Nf)^2 - \sum_{\mathbb{C}} s(\mathbb{C}) \sum_{\rho, \sigma=1}^{s(\mathbb{C})} |\tilde{\alpha}_{\rho\sigma}(\mathbb{C})|^2 = 0.$$

If we take all finite subsets of C for the $\mathbb{C}_1, \dots, \mathbb{C}_n$ in Theorem 26, we see that the left side of the above equation is the greatest lower bound of all $(N[f-g])^2$ if g is any linear aggregate of any finite number of elements $D_{\rho\sigma}(\mathbb{C})$. Hence we need to show merely that it can be made $\leq \epsilon^2$ for any $\epsilon > 0$. This is accomplished by choosing the unit $\phi(x)$ according to Theorem 16, be-

cause $\phi \times f$ is a linear aggregate of a finite number of elements $D_{\rho\sigma}(\mathbb{G})$. By Theorem 23 we need to prove this only for the elements $D_{\tau\nu}(\mathbb{D}) \times f$ and it follows that

$$\begin{aligned} D_{\tau\nu}(\mathbb{D}) \times f(x) &= M_{\nu}[D_{\tau\nu}(xy^{-1}; \mathbb{D})f(y)] \\ &= \sum_{\lambda=1}^{s(\mathbb{D})} M_{\nu}[D_{\tau\lambda}(x; \mathbb{D})D_{\lambda\nu}(y^{-1}; \mathbb{D})f(y)] \\ &= \sum_{\lambda=1}^{s(\mathbb{D})} M_{\nu}[D_{\lambda\nu}(y^{-1}; \mathbb{D})f(y)]D_{\tau\lambda}(x; \mathbb{D}). \end{aligned}$$

THEOREM 29. (*Approximation Theorem.*) *For every a.p. function $f(x)$ and every $\epsilon > 0$ there exists a linear aggregate $h(x)$ of a finite number of elements $D_{\rho\sigma}(x; \mathbb{G})$ (which can be limited to such elements for which the expansion matrix $\tilde{\alpha}(\mathbb{G})$ of f is $\neq 0$) such that $|f(x) - h(x)| \leq \epsilon$ for every x .*

By Theorem 18 we may put $h = \phi \times f \times g$, so that we need to prove merely that $\phi \times f \times g$ is a linear aggregate of the desired kind. By Theorem 23 we may consider $D_{\tau\nu}(\mathbb{D}) \times f \times g$. The last formula of the preceding proof shows that this is a linear aggregate of a finite number of elements $D_{\tau\lambda}(\mathbb{D})$. This formula gives for the coefficient of $D_{\tau\lambda}(\mathbb{D})$ in $D_{\rho\sigma}(\mathbb{D}) \times f \times g$ (replace its f by $f \times g$)

$$\begin{aligned} M_{\nu}[D_{\lambda\nu}(y^{-1}; \mathbb{D})f \times g(y)] &= M_{\nu}[\overline{D_{\lambda\nu}(y; \mathbb{D})'}f \times g(y)] \\ &= M_{\nu}[f \times g(y)\overline{D_{\nu\lambda}(y; \mathbb{D})}]. \end{aligned}$$

Thus it is the expansion coefficient of $D_{\nu\lambda}(\mathbb{D})$ in $f \times g$, and this is equal to zero if the expansion matrix of $D_{\nu\lambda}(\mathbb{D})$ in f is zero (cf. the statement of Theorem 25 concerning $f \times g$).

THEOREM 30. *Each a.p. function is the limit of a uniformly convergent sequence of functions each of which is a linear aggregate of a finite number of elements $D_{\rho\sigma}(\mathbb{G})$, and conversely.*

The statement follows from Theorem 29 by putting $\epsilon = 1, \frac{1}{2}, \frac{1}{3}, \dots$ in succession. The converse statement is a consequence of the a.p. character of all elements $D_{\rho\sigma}(\mathbb{G})$.

IV. ALMOST PERIODICITY AND CLOSED FAMILIES OF FUNCTIONS

14. Parts I–III give a fairly complete theory of a.p. functions in an arbitrary group $\mathbb{G}$, absolutely free from the customary restriction of continuity. We now introduce restrictions of this type, but in a more general manner, by considering certain families of functions.

DEFINITION 13. A set S of functions $f(x)$ (defined in $\mathfrak{G}$ with complex numbers as values) is called a closed family (cl.f.) if it has the following properties:

- (1) If $f(x)$ is in S , every $f(xa)$ is in S .
- (2) If $f(x)$ is in S , every $f(ax)$ is in S .
- (3) If $f(x)$ is in S , every $\alpha f(x)$ is in S .
- (4) If $f(x)$ and $g(x)$ are in S , $f(x) \pm g(x)$ is in S .
- (5) If $f_1(x), f_2(x), \dots$ are in S and if $f_n(x)$ converges uniformly to $f(x)$ as $n \rightarrow \infty$, then $f(x)$ is in S .

THEOREM 31. If S is a cl.f. and contains either f or g , then it contains $f \times g$; if $D(a) = \{D_{\rho\sigma}(a)\}$ is an irreducible normal representation, S contains every $D_{\rho\sigma}$ if it contains one $D_{\rho\sigma}$; if the system of representative irreducible normal representations $D(a; \mathfrak{G})$ is given (cf. Definition 11), and if S contains f , then S contains all elements $D_{\rho\sigma}(\mathfrak{G})$ of every $D(\mathfrak{G})$ whose expansion matrix $\tilde{\alpha}(\mathfrak{G})$ (cf. Definition 12) is $\neq 0$.

Therefore Theorems 28 (Parseval's equation), 29 (Approximation Theorem), and 30 remain true if we restrict ourselves throughout to functions in S .

If f belongs to S , $f \times g$ belongs to S by the Remark following Definition 6. The case where g belongs to S can be reduced to the case where f belongs to S by replacing ab by ba in $\mathfrak{G}$. If a $D_{\rho\sigma}$ belongs to S , every $D_{\rho'\sigma'}$ belongs to it, since, by Theorem 21, $s^2 D_{\rho'\rho} \times D_{\rho\sigma} \times D_{\sigma\sigma'} = D_{\rho'\sigma'}$. Finally,

$$\begin{aligned} M_y[f(y)D_{\sigma'\sigma}(y^{-1}x; \mathfrak{G})] &= M_y\left[f(y)\sum_{\tau=1}^s D_{\sigma'\tau}(y^{-1}; \mathfrak{G})D_{\tau\sigma}(x; \mathfrak{G})\right] \\ &= \sum_{\tau=1}^s M_y[f(y)D_{\sigma'\tau}(y^{-1}; \mathfrak{G})]D_{\tau\sigma}(x; \mathfrak{G}) = \sum_{\tau=1}^s M_y[f(y)\overline{D_{\tau\sigma'}(y; \mathfrak{G})}]D_{\tau\sigma}(x; \mathfrak{G}), \end{aligned}$$

that is,

$$\begin{aligned} f \times D_{\sigma'\sigma}(\mathfrak{G}) &= \sum_{\tau=1}^s \tilde{\alpha}_{\tau\sigma'}(\mathfrak{G})D_{\tau\sigma}(\mathfrak{G}), \\ D_{\rho\rho'}(\mathfrak{G}) \times f \times D_{\sigma'\sigma}(\mathfrak{G}) &= \frac{1}{s} \tilde{\alpha}_{\rho'\sigma'}(\mathfrak{G})D_{\rho\sigma}(\mathfrak{G}). \end{aligned}$$

Hence if $\tilde{\alpha}(\mathfrak{G}) \neq 0$ for a given $\mathfrak{G}$, that is, if any $\tilde{\alpha}_{\rho'\sigma'}(\mathfrak{G}) \neq 0$, and if f belongs to S , then $f \times D_{\sigma'\sigma}(\mathfrak{G})$, $D_{\rho\rho'}(\mathfrak{G}) \times f \times D_{\sigma'\sigma}(\mathfrak{G})$, and $D_{\rho\sigma}(\mathfrak{G})$ in turn belong to S .

If we keep these facts in mind we see that the proofs of Theorems 28, 29, and 30 still hold in S .

DEFINITION 14. *If a topology T is given in $\mathfrak{G}^\dagger$ we denote the set of all T -continuous functions by $[T]$.*

THEOREM 32. *If a topology T is given in $\mathfrak{G}^\dagger$ in which ab is continuous in a for a fixed b , and in b for a fixed a , then $[T]$ is a c.f.f.*

The statement is obvious.

The a.p. functions of a c.f.f. S are determined (in the manner described in Theorem 30; cf. the last statement of Theorem 31) by the elements $D_{\rho\sigma}(\mathfrak{G})$ belonging to it. This greatly facilitates the determination of all a.p. functions of a given c.f.f. S .

15. We shall discuss some examples in detail.

EXAMPLE 1. Let $\mathfrak{G} = \mathfrak{G}_{\text{rat}}$ be the set of all rational numbers with addition as the rule of composition. As the group is Abelian, all irreducible representations are of degree 1, $D(a; \mathfrak{G}) = \{D_{\rho\sigma}(a; \mathfrak{G})\}$, $\rho, \sigma = 1$, so that we have a single element $D_{11}(a; \mathfrak{G}) = \phi(a; \mathfrak{G})$. The fact that this is a unitary representation is expressed by the relation

$$(*) \quad \phi(a)\phi(b) = \phi(a+b), \quad |\phi(a)| = 1.$$

Every rational number can be written in the form $a = m/n!$ ($m = 0, \pm 1, \pm 2, \dots$; $n = 1, 2, \dots$). Now put

$$\phi\left(\frac{1}{n!}\right) = e^{2\pi\lambda_n i}, \quad 0 \leq \lambda_n < 1.$$

Then

$$(**) \quad \phi(a) = \phi\left(\frac{m}{n!}\right) = e^{2\pi\lambda_n m i}.$$

Since

$$\frac{n+1}{(n+1)!} = \frac{1}{n!},$$

it follows that

$$(*)_{**} \quad (n+1)\lambda_{n+1} \equiv \lambda_n \pmod{1} \quad (n = 1, 2, \dots)$$

and, on the other hand, it is clear that $(*)_{**}$ makes $(**)$ a definition prescribing a unique value for $\phi(a)$ which satisfies $(*)$. So the general solution is $(**)$, with the further condition $(*)_{**}$. An alternative way of writing $(*)_{**}$ is

$$(**)_{**} \quad \lambda_{n+1} \equiv \frac{\lambda_n + p_n}{n+1}, \quad p_n = 0, 1, \dots, n; \quad n = 1, 2, \dots$$

† See first footnote on page 456.

EXAMPLE 2. Take the same $\mathfrak{G} = \mathfrak{G}_{\text{rat}}$, but take its normal topology $T = T_0$ (distance $|a-b|$) and consider $S = [T_0]$. The question then is, for which λ_n 's does the $\phi(a)$ of (**) belong to $[T_0]$? That is, when is it T_0 -continuous? It is obvious that this means that $n!\lambda_n$ is bounded, and as (**) implies that $n!\lambda_n = \lambda_1 + 1!p_1 + \cdots + (n-1)!p_{n-1}$, it means that only a finite number of the p_m 's are $\neq 0$. Thus $n!\lambda_n$ is ultimately constant, say λ , and we have

$$(\S) \quad \phi(a) = e^{2\pi\lambda a} \quad (\lambda \text{ real}). \dagger$$

EXAMPLE 3. Take the same $\mathfrak{G} = \mathfrak{G}_{\text{rat}}$, but take its p -adic topology $T = T_p$ ($p = 2, 3, 5, \dots$ a prime number; distance is then 2^{N_0} , where N_0 is the minimal exponent $N = 0, \pm 1, \pm 2, \dots$ for which the least denominator of $p^N(a-b)$ is not divisible by p) and consider $S = [T_p]$. The question is, for which λ_n 's is the $\phi(a)$ of (**) T_p -continuous? In T_p , $p^n/n! \rightarrow 0$ as $n \rightarrow \infty$ ($n = 1, 2, \dots$, but fixed), so that $\exp(2\pi\lambda_n p^n i) \rightarrow 1$, $\lambda_n p^n \rightarrow 0 \pmod{1}$, which implies, of course, that there is a $\nu = \nu_n$ for which $\lambda_n p^\nu$ is an integer. This can be expressed in the following manner: there is a ν for which $\lambda_1 p^\nu$ is an integer, and p_n in (**) is divisible by the greatest divisor of $n+1$ which is prime to p . On the other hand, it is not difficult to see that this condition is sufficient.

EXAMPLE 4. Let $\mathfrak{G} = \mathfrak{G}_{\text{real}}$ be the set of all real numbers with addition as the rule of composition. Again we first determine all a.p. functions, that is, all irreducible unitary representations. This again means solving (*), but now with a and b running over all real numbers. Equation (*) can be solved by the following procedure:

Choose a rational linear basis of the set of real numbers, that is, a set B such that for every real number a the equation $a = \alpha_1 \xi_1 + \cdots + \alpha_n \xi_n$ ($n = 1, 2, \dots$; $\alpha_1, \dots, \alpha_n$ rational numbers, all $\neq 0$; $\xi_1, \dots, \xi_n$ different elements of B) has exactly one solution [12, pp. 459–462]. For every ξ of B we can define the quantity

$$\Gamma^{(a)}(\xi) = \begin{cases} \alpha_m & \text{if } \xi = \text{some } \xi_m, \\ 0 & \text{if } \xi \neq \text{each } \xi_m, \end{cases} \quad m = 1, \dots, n;$$

then $\Gamma^{(a)}(\xi)$ is always rational, we have

$$a = \sum_{\xi \in B} \Gamma^{(a)}(\xi) \xi,$$

where only a finite number of terms are $\neq 0$, and thus $\Gamma^{(a+b)}(\xi) = \Gamma^{(a)}(\xi) + \Gamma^{(b)}(\xi)$. From this it follows at once that every solution $\phi(a)$ of (*) for real a 's is of the form

[†] Thus there exist discontinuous a.p. functions of a rational variable. This fact was proved by Ursell [28, Second Note].

$$(\dagger) \quad \phi(a) = \prod_{\xi \in B} \phi_{\xi}(\Gamma^{(a)}(\xi)),$$

where each $\phi_{\xi}(c)$ is a solution of $(*)$ for rational c 's, and thus only a finite number of factors are $\neq 1$. Conversely, it is obvious that every $\phi(a)$ in $(\dagger)$ is a solution. Therefore the general solution is given by $(\dagger)$ if, for every ξ of B , we choose a $\phi_{\xi}(c)$ from $(**)$ and $(***)$ with $\lambda_n = \lambda_{\xi, n}$ and $p_n = p_{\xi, n}$ dependent on ξ .

EXAMPLE 5. Take the same $\mathfrak{G} = \mathfrak{G}_{\text{real}}$, but consider the set of all Lebesgue-measurable functions, $S = S_m$, which is obviously a c.l.f. The question is, which functions $\phi(a)$ of $(\dagger)$ are Lebesgue-measurable? As they are solutions of $\phi(a)\phi(b) = \phi(a+b)$ and $|\phi(a)| = 1$, we can infer from their measurability that they must be of the form

$$(\S) \quad \phi(a) = e^{2\pi\lambda a i} \quad (\lambda \text{ real}). \dagger$$

EXAMPLE 6. Take the same $\mathfrak{G} = \mathfrak{G}_{\text{real}}$, but take its normal topology $T = T_0$ (distance $|a-b|$) and consider $S = [T_0]$. The question is, which functions $\phi(a)$ of $(\dagger)$ are T_0 -continuous? As every T_0 -continuous function is measurable, all such functions must be of the form $(\S)$; and as all functions $(\S)$ are continuous, this again gives the general solution.||

EXAMPLE 7. Take the same $\mathfrak{G} = \mathfrak{G}_{\text{real}}$, but in it take a new topology $T = T(\lambda_1, \dots, \lambda_k)$, where the only relation $n_1\lambda_1 + \dots + n_k\lambda_k = 0$, with $n_1, \dots, n_k = 0, \pm 1, \pm 2, \dots$, shall be $n_1 = \dots = n_k = 0$; distance is defined by¶

$$\begin{aligned} & [|e^{2\pi\lambda_1 a i} - e^{2\pi\lambda_1 b i}|^2 + \dots + |e^{2\pi\lambda_k a i} - e^{2\pi\lambda_k b i}|^2]^{1/2} \\ &= 2[\sin^2 \pi\lambda_1(a-b) + \dots + \sin^2 \pi\lambda_k(a-b)]^{1/2}. \end{aligned}$$

Hence the condition $a_n \rightarrow a$ in $T(\lambda_1, \dots, \lambda_k)$ as $n \rightarrow \infty$ means that $a_n \rightarrow a$ with

† This is analogous to a result of Fréchet [9] who discussed $f(a)+f(b)=f(a+b)$. Cf. also Sierpinski [23] and Banach [1]. The simplest way to prove our statement is this:

Put $\psi_{\epsilon}(a) = \int_a^{a+\epsilon} \phi(x) dx$. Then $\psi_{\epsilon}(a)$ is continuous in a and satisfies $\psi_{\epsilon}(a)\phi(b) = \psi_{\epsilon}(a+b)$. If we had $\psi_{\epsilon}(a) = 0$ for every ϵ , then, as $(\partial/\partial \epsilon)\psi_{\epsilon}(a)$ is equal to $\phi(a+\epsilon)$ except over a set of measure zero, it would lead to a function $\phi(a+\epsilon) = 0$, except over a set of measure zero, which contradicts the condition $|\phi(a+\epsilon)| = 1$. Thus we can find ϵ_0 and a_0 such that $\psi_{\epsilon_0}(a_0) \neq 0$, and then our equation shows that $\phi(b) = \psi_{\epsilon_0}(a+b)/\psi_{\epsilon_0}(a_0)$, that is, $\phi(b)$ is continuous. Then $\psi_{\epsilon_0}(a)$ is differentiable, so that (by our last equation) $\phi(b)$ is also. Now we differentiate $\phi(a)\phi(b) = \phi(a+b)$ and get $\phi'(a)\phi(b) = \phi'(a+b)$, that is, $\phi'(b) = \beta\phi(b)$ when $a=0$. This means that $\phi(a) = \alpha e^{\beta a}$, and our original conditions make $\alpha = 1$, $\beta = 2\pi\lambda i$, λ real.

|| Thus there exist discontinuous a.p. functions of a real variable, but they are all non-measurable. These facts have also been proved by Ursell [28, First Note].

¶ For $k=1$ this is not only a new topology in $\mathfrak{G}_{\text{real}}$, but this also implies an identification of elements congruent mod $1/\lambda_1$. After this identification it is the normal topology. For $k>1$ it implies no identifications, but it is a new topology.

respect to mod $1/\lambda_1, \dots$, and mod $1/\lambda_k$ simultaneously. Therefore $\mathfrak{G}_{\text{real}}$ is compact when metrically completed in this topology and every uniformly $T(\lambda_1, \dots, \lambda_k)$ -continuous function is a.p. (cf. Theorem 36). The question is, which functions $\phi(a)$ of (†) are uniformly $T(\lambda_1, \dots, \lambda_k)$ -continuous? As $T(\lambda_1, \dots, \lambda_k)$ -continuity implies T_0 -continuity, they must have the form (§), that is, $\phi(a) = e^{2\pi\lambda a i}$. Now the condition $a_n \rightarrow a$ with respect to mod $1/\lambda_1, \dots$, and mod $1/\lambda_k$ should imply that $\phi(a_n) \rightarrow \phi(a)$ so that $e^{2\pi\lambda a_n i} \rightarrow e^{2\pi\lambda a i}$. This is the case if and only if

$$\lambda = n_1\lambda_1 + \dots + n_k\lambda_k, \quad n_1, \dots, n_k = 0, \pm 1, \pm 2, \dots.*$$

EXAMPLE 8. Let $\mathfrak{G}$ be a semi-simple Lie group.† The determination of all a.p. functions again means the determination of all irreducible unitary representations (which now of course need not be of degree 1). But such a representation is always continuous in the normal topology T_0 of $\mathfrak{G}$.‡ Therefore all a.p. functions in this $\mathfrak{G}$ are automatically T_0 -continuous, in contrast with Examples 1, 2, and 4, 6 (cf. footnotes † and || on pages 478 and 479 respectively). Thus there is no need to discuss $S = [T_0]$ separately.

16. Examples 1–8 sufficiently illustrate the various possibilities of combining a.p. functions with topology to make further comment unnecessary. We shall now investigate another phenomenon.

THEOREM 33. *Let $\mathfrak{G}$ be a group and S a cl.f. of functions in it. The following conditions on two elements a and b of $\mathfrak{G}$ are equivalent:*

- A. $D(a; \mathfrak{C}) = D(b; \mathfrak{C})$ (that is, $D_{\rho\rho}(a; \mathfrak{C}) = D_{\rho\rho}(b; \mathfrak{C})$) for every $\mathfrak{C}$ of C for which the elements $D_{\rho\rho}(\mathfrak{C})$ belong to S .
- B. $D(a) = D(b)$ (that is, $D_{\rho\rho}(a) = D_{\rho\rho}(b)$) for every normal representation for which the elements $D_{\rho\rho}$ belong to S .
- C. $f(a) = f(b)$ for every a.p. function in S .

A is a special case of B, so that B implies A.

As every $D_{\rho\rho}(x)$ is a.p. (Theorem 19 and Definition 10), B is a special case of C so that C implies B.

Finally, Theorems 30 and 31 show that A implies C.

Our three statements together prove the equivalence of A, B, and C.

DEFINITION 15. *We call two elements a and b of $\mathfrak{G}$ which satisfy one of the equivalent conditions of Theorem 33 S -coherent (if S is the set of all functions, we abbreviate this to coherent). We denote the set of those elements which are S -coherent (coherent) with 1 by $\mathfrak{G}^S$ ($\mathfrak{G}_0$).*

* That is, we are led to the Bohl-Esclangon [3] quasi-periodic functions with the basis $\lambda_1, \dots, \lambda_n$. Cf. also H. Bohr [4, II, pp. 111–117].

† For a detailed discussion of this notion cf. E. Cartan [7].

‡ This is a most remarkable difference between the behavior of Abelian Lie groups (cf. Example 4) and semi-simple Lie groups. It was discovered by B. L. van der Waerden [29, p. 785].

THEOREM 34. $\mathfrak{G}^S$ is an invariant subgroup of $\mathfrak{G}$, and if $S = [T]$ for a topology T of $\mathfrak{G}$, then $\mathfrak{G}^S$ is T -closed. Those elements of $\mathfrak{G}$ which are coherent with a given a form the coset of $\mathfrak{G}_0^S$ in $\mathfrak{G}$ belonging to a .

Consider the condition B in Theorem 33 (either A or C could also be used). If a and b belong to $\mathfrak{G}_0^S$ we have $D(ab) = D(a)D(b) = 1$, $D(a^{-1}) = D(a)^{-1} = 1$, that is, a^{-1} and ab belong to it; if only a belongs to $\mathfrak{G}_0^S$, we have $D(b^{-1}ab) = D(b)^{-1}D(a)D(b) = D(b)^{-1}D(b) = 1$, that is, $b^{-1}ab$ also belongs to it. If $S = [T]$, every $D(a)$ is T -continuous, each set $D(a) = 1$ is T -closed, and so their common part $\mathfrak{G}_0^S$ is also. That a and b are coherent means that we always have $D(a) = D(b)$, $D(a^{-1}b) = D(a)^{-1}D(b) = 1$, that is, that $a^{-1}b$ belongs to $\mathfrak{G}_0^S$. Hence the elements b form exactly the coset of $\mathfrak{G}_0^S$ in $\mathfrak{G}$ belonging to a .

DEFINITION 16. If $\mathfrak{G}_0^S = 1$ ($\mathfrak{G}_0 = 1$) we call $\mathfrak{G}$ and S ($\mathfrak{G}$) *maximally a.p.*; if $\mathfrak{G}_0^S = \mathfrak{G}$ ($\mathfrak{G}_0 = \mathfrak{G}$) we call $\mathfrak{G}$ and S ($\mathfrak{G}$) *minimally a.p.*

These two cases are indeed the two extremes which can occur. If $\mathfrak{G}$ and S are minimally a.p., then for every a.p. $f(x)$ of S we always have $f(a) = f(1)$, that is, the constants are the only a.p. functions in S . And for every normal representation $D(a)$ with the elements $D_{\rho\sigma}(a)$ in S , it must be $D(a) = D(1) = 1$, so that if $D(a)$ is irreducible its degree must be 1. If, on the other hand, $\mathfrak{G}$ and S are maximally a.p., then there exists, for every pair a and b in $\mathfrak{G}$, $a \neq b$, a $\mathfrak{C}$ from C such that all $D_{\rho\sigma}(\mathfrak{C})$ are in S with $D(a; \mathfrak{C}) \neq D(b; \mathfrak{C})$, and an a.p. function $f(x)$ in S such that $f(a) \neq f(b)$. Even more is true:

THEOREM 35. If $f(x)g(x)$ is in S whenever $f(x)$ and $g(x)$ are in S , and if $\mathfrak{G}$ and S are maximally a.p., then, for any finite set $a_1, \dots, a_n$ of distinct elements of $\mathfrak{G}$ and any set of complex numbers $\alpha_1, \dots, \alpha_n$, an a.p. function $f(x)$ exists in S with the prescribed values $f(a_1) = \alpha_1, \dots, f(a_n) = \alpha_n$.

If $a \neq b$, there is an a.p. function $g(x)$ in S such that $g(a) \neq g(b)$, so that

$$h(x) = \frac{g(x) - g(b)}{g(a) - g(b)}$$

is an a.p. function in S with $h(a) = 1$ and $h(b) = 0$. For every pair a and b ($a \neq b$), choose such a function $h(x)$ and denote it by $h(a, b; x)$. Then

$$f(x) = \sum_{\nu=1}^n \alpha_{\nu} \prod_{\mu=1, \mu \neq \nu}^n h(a_{\mu}, b_{\mu}; x)$$

has all the properties required.

There are also some other ways to characterize $\mathfrak{G}_0^S$, but we shall not discuss them here.

17. If $\mathfrak{G}$ and S are maximally a.p., we can introduce a topology by means of their a.p. functions. In this connection the following notions are of importance:

DEFINITION 17. If $\mathfrak{G}$ and S are maximally a.p. we define a topology FS in $\mathfrak{G}$ by considering the following "neighborhoods" $\mathfrak{N}(a)$ of an element a of $\mathfrak{G}$ †: Choose a finite number of a.p. functions $f_1, \dots, f_n$ and an $\epsilon > 0$; then $\mathfrak{N}(a) = \mathfrak{N}(a; f_1, \dots, f_n, \epsilon)$ is the set of all b 's such that $|f_1(a) - f_1(b)| < \epsilon, \dots, |f_n(a) - f_n(b)| < \epsilon$. (If S is the set of all functions we abbreviate this to F .)

One sees at once that FS satisfies Hausdorff's Axioms (cf. first footnote on page 447).

DEFINITION 18. If two topologies T_1 and T_2 for a set $\mathfrak{S}$ are given, T_1 is called weaker than T_2 if every T_1 -neighborhood of an element a of $\mathfrak{S}$ contains a T_2 -neighborhood of a .

Obviously, every set which is closed or open in the T_1 -sense, and every function which is continuous in the T_1 -sense, has the same property in the T_2 -sense. Thus for a group $\mathfrak{S} = \mathfrak{G}$, $[T_1]$ is a subset of $[T_2]$ (cf. Definition 14). On the other hand, it is obvious that if S_1 and S_2 are cl.f. and S_2 is a subset of S_1 , then FS_1 is weaker than FS_2 .

We intend to go more deeply into the theory of $[T]$ and FS on another occasion. At present let us merely remark that for every $\mathfrak{G}$ (even for a non-topologically given $\mathfrak{G}$) F is a topology determined by $\mathfrak{G}$ alone (if $\mathfrak{G}$ is maximally a.p.). Discussion of Examples 1 and 4 shows without much difficulty that $\mathfrak{G}_{\text{rat}}$ and $\mathfrak{G}_{\text{real}}$ are maximally a.p. (even with their cl.f. $[T_0]$ or $[T_p]$ and $[T_0]$ or $[T(\lambda_1, \dots, \lambda_k)]$ respectively (for $k > 1$, cf. footnote* on page 480)) and that their F 's are very "strong"; the condition $a_n \rightarrow a$ in F as $n \rightarrow \infty$ means that all a_n 's, with a finite number of exceptions, are equal to a . On the other hand, Example 8 shows that, for a semi-simple Lie group $\mathfrak{G}$ (if it is maximally a.p.), $F = F[T_0]$. Theorem 36 shows that if $\mathfrak{G}$ is compact in T_0 , $\mathfrak{G}$ and $[T_0]$ are maximally a.p. and $F[T_0] = T_0$. Thus, if $\mathfrak{G}$ is a semi-simple and compact Lie group, it is maximally a.p. and $F = T_0$.

18. The case where a group $\mathfrak{G}$ and a topology T have the properties that $\mathfrak{G}$ and $[T]$ are maximally a.p. and $F[T] = T$ is of particular importance.

THEOREM 36. $\mathfrak{G}$ and $[T]$ are maximally a.p. and $F[T] = T$ in each of the two following cases (in case B, $F[T] = T$ should be understood to mean only that the condition $a_n \rightarrow a$ as $n \rightarrow \infty$ is equivalent to it in both senses):

A. $\mathfrak{G}$ is compact in T .

† See first footnote on page 456

B. $\mathfrak{G}$ is locally compact[†] and separable[†] in T , $\mathfrak{G}$ is an Abelian group, and ab and a^{-1} are T -continuous in a and b , and a respectively.[‡]

If $\mathfrak{G}$ is compact in T , every continuous $f(x)$ is a.p.: for $\mathfrak{G}$ being compact, $f(x)$ is uniformly continuous; if any sequence $a_1, a_2, \dots$ is given, we can extract from it a subsequence $a_{n_1}, a_{n_2}, \dots$ which converges to a limit a , and then we have the result that $f(xa_{n_i}) \rightarrow f(xa)$ and $f(a_{n_i}x) \rightarrow f(ax)$ uniformly as $i \rightarrow \infty$. Thus $[T]$ consists only of a.p. functions.

Now it is possible to define a distance $D(a, b)$ in $\mathfrak{G}$ which is equivalent to the topology T [27, 25]. $f(x) = D(a, x)$ belongs to $[T]$ and we have the result that $f(a) = 0 \neq f(b)$, proving that $\mathfrak{G}$ and $[T]$ are maximally a.p.; and the neighborhood $\mathfrak{N}(a; f, \epsilon)$ (cf. Definition 17) is the sphere with the center a and the radius ϵ , proving that T is weaker than $F[T]$. But $F[T]$ is obviously weaker than T , and therefore $F[T] = T$. Thus A is proved.

The proof of B will be given at the end of Part V.

Minimally a.p. groups likewise exist, for example, the group $g_{(n)}$ of all linear transformations of determinant 1 in the real euclidean space of n dimensions, $n=2, 3, \dots$. As it is a semi-simple Lie group, indeed even simple [6], all its bounded representations are continuous [29]; as it is a linear group, their continuity implies their differentiability [14, p. 37]. Hence we need only determine those irreducible representations of $g_{(n)}$ which arise from "infinitesimal representations," and see if there exist any bounded ones among them. Now these representations and their traces (characteristics) are known [70; 30, pp. 287, 300], and only the identity, $D(a) \equiv 1$, has a bounded trace, so no other representation can be bounded. Application of Theorem 33, criterion A, shows that $g_{(n)}$ is minimally a.p.

The group g' of all transformations $y = ax + b$ (a and b real, $a > 0$) is neither minimally nor maximally a.p., as a simple discussion shows.

V. ABELIAN GROUPS

19. We assume throughout Part V that the assumptions of Theorem 36, case B (which we shall finally prove), hold; thus we assume that a group $\mathfrak{G}$ and a topology T are given, that $\mathfrak{G}$ is locally compact and separable in T and Abelian, and that ab, a^{-1} are T -continuous.

Under the above topological assumptions, A. Haar has shown the existence of a right-invariant Lebesgue integral [11, pp. 166–167]. Thus it is possible to define for complex-valued functions $f(x)$ defined in $\mathfrak{G}$ (i) a notion of measura-

[†] "Locally compact" means that each element a has a conditionally compact neighborhood [13, p. 107]; as a group $\mathfrak{G}$ is homogeneous it is sufficient to postulate this for the element 1. "Separable" means that there exists a countably infinite "equivalent system of neighborhoods"; if the topology is originated by a distance notion, one may postulate the existence of a countably infinite everywhere dense subset [13, p. 125, and p. 229, Axiom 10].

[‡] See first footnote on page 456.

bility, (ii) a notion of summability, (iii) an integral $\int_{\mathfrak{G}} f(x) dx$. On the basis of (i) and (ii), moreover, it is possible to do this in such a manner that (i)–(iii) have all the formal properties of these notions as in the usual Lebesgue theory, and besides are invariant under the substitution of $f(xa)$ for $f(x)$.

We now consider all measurable functions $f(x)$ in $\mathfrak{G}$ for which $|f(x)|^2$ is summable, that is, $\int_{\mathfrak{G}} |f(x)|^2 dx$ is finite. These functions form a Hilbert space $\mathfrak{H}$ if we define the inner product (f, g) to be $\int_{\mathfrak{G}} f(x) \overline{g(x)} dx$ †, provided that $\mathfrak{G}$ is infinite, which we will assume to be the case. (If it is finite, it is compact and falls under case A.) In $\mathfrak{H}$,

$$(\#) \quad O_a f(x) = f(xa)$$

defines a linear and unitary operation (that is, an operation which leaves (f, g) invariant), and it follows that

$$(\#\#) \quad O_a O_b = O_{ab}.$$

Now we use the Abelian character of $\mathfrak{G}$, by virtue of which $(\#\#)$ implies that O_a and O_b commute. As O_a is unitary, its adjoint† is $O_a^* = O_a^{-1}$ and, by $(\#\#)$, $O_a^* = O_{a^{-1}}$. Thus every O_b commutes with every O_a and O_a^* , and the set of all operators O_a has been called Abelian [16, p. 389]. Therefore a theorem proved by the author applies to this set: there exists a bounded Hermitian operator R such that every O_a is a function of R ,

$$(\#\#\#) \quad O_a = \phi_a(R),$$

where $\phi_a(\lambda)$ is a complex-valued function of the variable λ .§ That the functions $\phi_a(\lambda)$ can be used for the discussion of the group $\mathfrak{G}$ has been noted by Haar and successfully applied to countably infinite Abelian groups [10, p. 131]; cf. also Wiener and Paley [33]. Theorem 37 will be an application of this idea in the full generality allowed by Haar's right-invariant Lebesgue integral. It must be remarked, however, that Haar's method of discussing countably infinite Abelian groups has been considerably simplified by Wiener and Paley [33], but that their simplification seems not to apply to our general case, and that we have to use Haar's original method.

THEOREM 37. *If $\mathfrak{G}$ and T fulfill the assumptions formulated at the beginning of this part (that is, if $\mathfrak{G}$ is locally compact, separable, and Abelian), there exists a function in two variables $\phi(a, \lambda)$ (a in $\mathfrak{G}$, λ real) with the following properties:*

† For the modern theory of Hilbert space cf. J. v. Neumann [15, pp. 63–70, 108–111]. Cf. further M. H. Stone [24, pp. 1–32].

§ The notion of a function of an operator is due originally to F. Riesz. More general forms have been given to it by J. v. Neumann [17, pp. 202–213] and M. H. Stone [24, pp. 221–241]. The theorem in question has been proved by J. v. Neumann [17, p. 214].

$\phi(a, \lambda)$ is a Baire function in (a, λ) ,* and there exists a "resolution of the identity" $E(\lambda)$ such that

$$(\dagger) \quad (O_a f, g) = \int_{-\infty}^{\infty} \phi(a, \lambda) d(E(\lambda)f, g) \dagger$$

identically in a and $f(x)$ and $g(x)$.

For the function $\phi(a, \lambda) = \phi_a(\lambda)$ in $(\frac{a}{\#}, \frac{\lambda}{\#})$, the theorem mentioned in the footnote§ on page 484 leads to all our statements except for the Baire character of $\phi(a, \lambda)$ in (a, λ) (it would show the Baire character in λ , but we need it also in a).

We know that finite linear aggregates of functions $f_{\mathfrak{D}}$ where $\mathfrak{D}$ is a conditionally compact open set, therefore having finite measure, and

$$f_{\mathfrak{D}}(a) = \begin{cases} 1 & \text{for } a \text{ in } \mathfrak{D}, \\ 0 & \text{elsewhere,} \end{cases}$$

are everywhere dense in our functional space [15, p. 110]. From now on "everywhere dense" will be interpreted in the sense of the distance

$$D(f, g) = \|f - g\| = [(f - g, f - g)]^{1/2} = \left[\int_{\mathfrak{G}} |f(x) - g(x)|^2 dx \right]^{1/2}$$

but not in the sense of the distance l.u.b._x $|f(x) - g(x)|$. Now, if $\mathfrak{D}_1$ is an open set the closure of which is part of $\mathfrak{D}$, we can find a continuous function§

$$f_{\mathfrak{D}, \mathfrak{D}_1}(a) \begin{cases} = 1 & \text{for } a \text{ in } \mathfrak{D}, \\ = 0 & \text{for } a \text{ not in } \mathfrak{D}, \\ \geq 0 \text{ and } \leq 1 & \text{elsewhere.} \end{cases}$$

If we let $\mathfrak{D}_1$ converge to $\mathfrak{D}$, then $f_{\mathfrak{D}, \mathfrak{D}_1}(a)$ converges everywhere to, and is majorized by, $f_{\mathfrak{D}}(a)$, so that $f_{\mathfrak{D}}(a)$ is its limit in the sense of the distance $\|f - g\|$. Therefore continuous functions f which are $\neq 0$ only in conditionally compact sets are everywhere dense in our functional space. Since $a_n \rightarrow a$ implies $xa_n \rightarrow xa$, we have $f(xa_n) \rightarrow f(xa)$ for these functions, and, by the second property,

$$\|O_{a_n} f - O_a f\| = \left[\int_{\mathfrak{G}} |f(xa_n) - f(xa)|^2 dx \right]^{1/2} \rightarrow 0.$$

Hence $a_n \rightarrow a$ implies $O_{a_n} f \rightarrow O_a f$ for an everywhere dense set of f 's, but as all

* That is, it can be obtained from continuous functions in (a, λ) by successive limiting processes wherein the limit is always taken of everywhere convergent sequences.

† $\int_{-\infty}^{\infty}$ is a Lebesgue-Stieltjes integral over λ . For an explanation of the terminology used, see [15, p. 92] or [24, p. 174].

§ This is a problem of Fréchet, first solved by Hahn. Cf. [26, Anhang III, p. 290].

O_a 's are unitary operators, and therefore uniformly continuous in f , the implication holds for every f . Consequently O_af is a continuous function in a for every fixed f .

A simple computation shows, after substituting $E(\mu)g$ in (§), at the end of Theorem 37 in place of g [cf. 17, p. 206],

$$(O_af, E(\mu)g) = \int_{-\infty}^{\mu} \phi(a, \lambda) d(E(\lambda)f, g).$$

Now choose a complete normalized orthogonal system $f_1, f_2, \dots$, put $f = g = f_n$, $n = 1, 2, \dots$, multiply by 2^{-n} , and add. The infinite series thus obtained in the left- and right-hand members converge uniformly since $(O_af, E(\mu)g)$ and $(E(\lambda)f, g)$ are both $\leq \|f\| \|g\|$ in absolute value. The result is

$$\sum_{n=1}^{\infty} 2^{-n} (O_af_n, E(\mu)f_n) = \int_{-\infty}^{\mu} \phi(a, \lambda) d \left[\sum_{n=1}^{\infty} 2^{-n} (E(\lambda)f_n, f_n) \right].$$

$(O_af_n, E(\mu)f_n)$ is continuous in a and continuous on the right in μ , $(E(\lambda)f_n, f_n)$ is continuous on the right in λ and monotonically increasing, and the same properties hold for the uniformly convergent sums

$$F(a, \mu) = \sum_{n=1}^{\infty} 2^{-n} (O_af_n, E(\mu)f_n), \quad G(\lambda) = \sum_{n=1}^{\infty} 2^{-n} (E(\lambda)f_n, f_n).$$

Thus $F(a, \mu)$ and $G(\lambda)$ are Baire functions, the latter is monotonically increasing, and

$$F(a, \mu) = \int_{-\infty}^{\mu} \phi(a, \lambda) dG(\lambda).$$

If we consider $G(\lambda)$ as the variable (instead of λ), then the well known theorem on the differentiability of integrals shows that†

$$\lim_{\delta, \epsilon \rightarrow 0^+} \frac{F(a, \mu + \delta) - F(a, \mu - \epsilon)}{G(\mu + \delta) - G(\mu - \epsilon)}$$

exists and equals $\phi(a, \mu)$ except, however, for a set of μ 's dependent on a whose $\xi = G(\mu)$ -image§ is a set (of real numbers) of Lebesgue measure zero. Now the function

$$\phi_1(a, \lambda) = \begin{cases} \lim_{\delta, \epsilon \rightarrow 0^+} \frac{F(a, \mu + \delta) - F(a, \mu - \epsilon)}{G(\mu + \delta) - G(\mu - \epsilon)} & \text{when this limit exists,} \\ 0 & \text{otherwise} \end{cases}$$

† Cf. [5, pp. 544-545]. Analogous results concerning "central derivatives" of F with respect to G are due to Daniell [8].

§ If $G(\mu)$ is discontinuous at $\mu = \mu_0$, the image of $\mu = \mu_0$ is supposed to be the whole jump-interval $G(\mu_0 - 0) \leq \xi \leq G(\mu_0 + 0)$.

is obviously a Baire function, and the set Σ of λ 's for which $\phi_1(a, \lambda) \neq \phi(a, \lambda)$ has a $\xi = G(\mu)$ -image of Lebesgue measure zero. Since $2^n G(\mu) - (E(\mu)f_n, f_n)$ is monotonically increasing, the $\xi = (E(\mu)f_n, f_n)$ -image of Σ is also of Lebesgue measure† zero, and therefore every $\xi = (E(\mu)f, f)$ -image is of Lebesgue measure zero [17, p. 213, the last remark of Part II].

Hence we may replace $\phi(a, \lambda)$ in (‡) by $\phi_1(a, \lambda)$, and this will not affect the validity of (‡) for $f = g$; now if we replace our f by $(f+g)/2$ and $(f-g)/2$ and subtract, we get the real part of the general (‡); if we replace f and g by if and g , we get its imaginary part, and prove it altogether. Thus $\phi_1(a, \lambda)$ meets all our requirements.

THEOREM 38. *Under the assumptions of Theorem 37, $\phi(a, \lambda)$ can even be chosen as a continuous function in a satisfying the equations*

$$\phi(ab, \lambda) = \phi(a, \lambda) \phi(b, \lambda), \quad |\phi(a, \lambda)| = 1.$$

A simple computation shows [17, p. 206] that

$$(O_a O_b f, g) = \int_{-\infty}^{+\infty} \phi(a, \lambda) \phi(b, \lambda) d(E(\lambda)f, g),$$

$$(O_a^* O_a f, g) = \int_{-\infty}^{+\infty} |\phi(a, \lambda)|^2 d(E(\lambda)f, g);$$

on the other hand,

$$(O_a b f, g) = \int_{-\infty}^{+\infty} \phi(ab, \lambda) d(E(\lambda)f, g), \quad (f, g) = \int_{-\infty}^{+\infty} d(E(\lambda)f, g).$$

Now $O_a O_b = O_{ab}$, $O_a^* O_a = 1$; therefore the right sides of our equations are equal. An analogous computation shows that if we substitute $E(\mu)g$ for g [17, p. 206] and subtract, we get

$$\int_{-\infty}^{\mu} (\phi(ab, \lambda) - \phi(a, \lambda) \phi(b, \lambda)) d(E(\lambda)f, g) = 0,$$

$$\int_{-\infty}^{\mu} (|\phi(a, \lambda)|^2 - 1) d(E(\lambda)f, g) = 0.$$

Putting $f = g$ shows that the equations of our Theorem hold except for a set of λ 's (depending on the pair a and b and on a respectively), the $\xi = (E(\lambda)f, f)$ -image of which has Lebesgue measure zero† (this condition holds for all f 's simultaneously). Returning to the complete normalized orthogonal

† The Lebesgue measure of the $\xi = H(\mu)$ -image of a set $\mathfrak{S}$ is $\int_{\mathfrak{S}} dH(\mu)$, and therefore, if it is 0 for $H(\mu)$, it will be 0 for every other function $K(\mu)$ for which $H(\mu) - K(\mu)$ is monotonically increasing [cf. 17, p. 198, rule d , and p. 199].

system $f_1, f_2, \dots$ in the proof of Theorem 37, and to the corresponding

$$G(\lambda) = \sum_{n=1}^{\infty} 2^{-n} (E(\lambda) f_n, f_n),$$

we see that also the $\xi = G(\lambda)$ -images have Lebesgue measure zero (this follows for each $\sum_{n=1}^N 2^{-n} (E(\lambda) f_n, f_n)$ from the integral formula in the footnote on page 487, and for

$$G(\lambda) = \sum_{n=1}^{\infty} 2^{-n} (E(\lambda) f_n, f_n)$$

from the fact that the difference is monotonically increasing, ≥ 0 , and $\leq \sum_{n=N+1}^{\infty} 2^{-n} = 1/2^N$.

For a -sets and b -sets (that is, subsets of $\mathfrak{G}$) we have a measure, namely, the Haar-Lebesgue measure. For λ -sets (that is, sets of real numbers) we shall consider the Lebesgue measure of the $\xi = G(\lambda)$ -image (cf. the footnote§ on page 486), and call it the λ -measure. All these measures have the formal properties of the Lebesgue measure.* We can, by the analogue of the process which leads from linear to plane measure [cf. 18, p. 588] use these measures to define measures with similar properties for (a, b) -sets, (a, λ) -sets, and (a, b, λ) -sets. If we use these defining processes, the theorem of Fubini holds for all combinations of the variables a, b, λ because its proof [5, pp. 622–628] applies unchanged.

As we are dealing with Baire functions, the (a, b, λ) - and (a, λ) -sets for which $\phi(ab, \lambda) \neq \phi(a, \lambda)\phi(b, \lambda)$ and $|\phi(a, \lambda)| \neq 1$ are Borel sets and therefore measurable. Hence Fubini's theorem can be applied to them; as for fixed a, b and λ respectively they give λ -sets of zero λ -measure, they are sets of zero (a, b, λ) -measure and (a, λ) -measure themselves. This again implies that if λ does not belong to a certain (fixed) set $\mathfrak{S}_1$ of zero λ -measure, and if a does not belong to a certain set $\mathfrak{S}_2^{(\lambda)}$ (depending on λ) of zero (Haar) measure, then we have the result that, if b does not belong to a certain set $\mathfrak{S}_3^{(\lambda, a)}$ (depending on λ and a) of zero (Haar) measure, then $\phi(ab, \lambda) = \phi(a, \lambda)\phi(b, \lambda)$, and, at any rate, $|\phi(a, \lambda)| = 1$.

Now choose a conditionally compact open set $\mathfrak{D}$. If we had, for a certain λ ,

$$\int_{\mathfrak{D}} \phi(x, \lambda) dx = 0$$

for every $\mathfrak{D}$, this would imply that

$$\int_{\mathfrak{M}} \phi(x, \lambda) dx = 0$$

* By this we mean that they satisfy Carathéodory's postulates I–V [5, pp. 238–239, 258].

for every measurable set $\mathfrak{M}$, and thus $\phi(x, \lambda) = 0$ except for an x -set of λ -measure zero. This contradicts the fact that we have $|\phi(x, \lambda)| = 1$ except for an x -set of λ -measure zero. Therefore choose $\mathfrak{D}$ such that

$$\int_{\mathfrak{D}} \phi(x, \lambda) dx \neq 0,$$

and denote the set of all (ax) 's (x in $\mathfrak{D}$) by $a\mathfrak{D}$.

Assume λ in $\mathfrak{S}_1$ and a in $\mathfrak{S}_2^{(\lambda)}$. Then we have $\phi(ax, \lambda) = \phi(a, \lambda)\phi(x, \lambda)$ if x is not in $\mathfrak{S}_3^{(\lambda, a)}$ and this implies that

$$\begin{aligned} \phi(a, \lambda) \int_{\mathfrak{D}} \phi(x, \lambda) dx &= \int_{\mathfrak{D}} \phi(ax, \lambda) dx = \int_{a\mathfrak{D}} \phi(y, \lambda) dy, \\ \phi(a, \lambda) &= \frac{\int_{a\mathfrak{D}} \phi(x, \lambda) dx}{\int_{\mathfrak{D}} \phi(x, \lambda) dx}. \end{aligned}$$

Now, by well known theorems on Lebesgue integrals, the numerator is continuous in a , the denominator is constant and $\neq 0$, and for this argument we need not restrict a to $\mathfrak{S}_2^{(\lambda)}$ (λ , of course, is in $\mathfrak{S}_1$). Hence we may define a continuous function $\phi_2(a, \lambda)$ by putting it equal to

$$\frac{\int_{a\mathfrak{D}} \phi(x, \lambda) dx}{\int_{\mathfrak{D}} \phi(x, \lambda) dx}$$

(λ in $\mathfrak{S}_1$); and we have $\phi_2(a, \lambda) = \phi(a, \lambda)$ if a is not in $\mathfrak{S}_3^{(\lambda)}$.

In this case, if b is not in $\mathfrak{S}_2^{(\lambda)}$ and ab is not in $\mathfrak{S}_2^{(\lambda)}$, we have $\phi_2(ab, \lambda) = \phi_2(a, \lambda)\phi_2(b, \lambda)$. But as we except only a b -set of zero (Haar) measure, this holds in an everywhere dense b -set, and thus, for reasons of continuity, for every b . So the above formula is true for every b , and $|\phi_2(a, \lambda)| = 1$ is true if a is not in $\mathfrak{S}_2^{(\lambda)}$. For the same continuity reasons, therefore, there are no a -exceptions at all. Thus (if λ is not in $\mathfrak{S}_1$) $\phi_2(a, \lambda)$ meets the requirements of our theorem if it can replace $\phi(a, \lambda)$.

By definition, $\phi_2(a, \lambda)$ is a Baire function in (a, λ) . Hence the (a, λ) -set for which $\phi(a, \lambda) \neq \phi_2(a, \lambda)$ is a Borel set and therefore measurable. Hence Fubini's theorem applies to it, and since for a fixed λ , except in a set of zero λ -measure ($\mathfrak{S}_1$), it gives an a -set of zero (Haar) measure ($\mathfrak{S}_2^{(\lambda)}$), it is a set of zero (a, λ) -measure itself. This again implies that if a does not belong to a certain (fixed) set $\mathfrak{S}_1'$ of zero (Haar) measure, then $\phi(a, \lambda) = \phi_2(a, \lambda)$ provided λ does not belong to a certain set $\mathfrak{S}_2'^{(a)}$ (depending on a) of zero λ -measure. If we change $\phi_2(a, \lambda)$ for the λ 's of $\mathfrak{S}_1$ into 1, we obtain its continuity in a and the equations of our Theorem for all λ 's without exceptions, and the state-

ment just now proved still holds if we replace $\mathfrak{S}_2'^{(a)}$ by $\mathfrak{S}_2'^{(a)} + \mathfrak{S}_1$ which is also a set of zero λ -measure.

If a does not belong to $\mathfrak{S}_1'$, we have $\phi(a, \lambda) = \phi_2(a, \lambda)$ except for a λ -set with zero λ -measure, that is, with a $\xi = G(\lambda)$ -image (cf. the footnote§ on page 486) of zero Lebesgue measure. This proves, as at the end of the proof of Theorem 37, that

$$(O_a f, g) = \int_{-\infty}^{+\infty} \phi_2(a, \lambda) d(E(\lambda)f, g)$$

identically in $f(x)$ and $g(x)$ if a does not belong to $\mathfrak{S}_1'$. Since $\mathfrak{S}_1'$ has zero (Haar) measure, the domain of validity in a is everywhere dense. But both sides are continuous functions of a : this was shown for the left side at the beginning of the proof of Theorem 37, and follows for the right side from the continuity of $\phi_2(a, \lambda)$ in a for all λ 's. Hence our equation holds for all a 's.

Thus $\phi_2(a, \lambda)$ meets all our requirements.

THEOREM 39. *If the assumptions of Theorems 37 and 38 are satisfied, and if ab and a^{-1} are continuous in (a, b) and in a respectively, then the condition $\phi(a, \lambda) = \phi(b, \lambda)$ for all λ 's is equivalent to the condition $a = b$, and the condition $\phi(a_n, \lambda) \rightarrow \phi(a, \lambda)$ as $n \rightarrow \infty$ for all λ 's is equivalent to the condition $a_n \rightarrow a$ as $n \rightarrow \infty$.*

The first statement follows from the second by putting $a_1 = a_2 = \dots = b$. The necessity of the criterion in the second statement is obvious, as all the functions $\phi(a, \lambda)$ are continuous in a . So the only thing we need to prove is its sufficiency.

Therefore suppose that $\phi(a_n, \lambda) \rightarrow \phi(a, \lambda)$ as $n \rightarrow \infty$ for all λ 's. Then (§) in Theorem 37 shows that $(O_{a_n} f, g) \rightarrow (O_a f, g)$ as $n \rightarrow \infty$ for any $f(x)$ and $g(x)$ of our functional space. Now let $\mathfrak{D}$ be a conditionally compact open set, and define

$$f_{\mathfrak{D}}(x) = \begin{cases} 1 & \text{for } x \text{ in } \mathfrak{D}, \\ 0 & \text{elsewhere.} \end{cases}$$

Put $f(x) = f_{\mathfrak{D}}(x)$ and $g(x) = f_{\mathfrak{D}}(ax)$. Then

$$(O_{a_n} f, g) = \int_{\mathfrak{G}} f(a_n x) \overline{f(ax)} dx,$$

$$(O_a f, g) = \int_{\mathfrak{G}} |f(ax)|^2 dx > 0,$$

and $(O_{a_n} f, g) \rightarrow (O_a f, g)$ implies that if n is sufficiently large, $(O_{a_n} f, g) \neq 0$ and therefore $f(a_n x) f(ax) \neq 0$. Hence there is an x for which $a_n x$ and ax both belong

to $\mathfrak{O}$, and as $a_n a^{-1} = (a_n x)(ax)^{-1}$, $a_n a^{-1}$ can be written in the form uv^{-1} , where u and v both belong to $\mathfrak{O}$. As every open set has conditionally compact subsets, this holds for every open $\mathfrak{O}$.

Now let $\mathfrak{N}$ be a neighborhood of a . Then we can find an open set $\mathfrak{O}$ for which every uv^{-1} , u and v in $\mathfrak{O}$, belongs to $\mathfrak{N}$. (Here is where the extra continuity assumptions are used.) Then our result shows that if n is sufficiently large, a_n belongs to $\mathfrak{N}$. This means that $a_n \rightarrow a$ as $n \rightarrow \infty$.

Theorems 37 and 38, combined with Theorem 19, show that each $\phi(a, \lambda)$, when considered as an a -function, is a.p. and belongs to $[T]$. Therefore Theorem 39 proves exactly the statements of Theorem 36, case B.

BIBLIOGRAPHY

1. S. Banach, *Sur l'équation fonctionnelle* $f(x+y)=f(x)+f(y)$, *Fundamenta Mathematicae*, vol. 1 (1920), pp. 123-124.
2. S. Bochner, *Beiträge zur Theorie der fastperiodischen Funktionen*, *Mathematische Annalen*, vol. 96 (1927), pp. 119-147.
3. P. Bohl, *Über die Darstellung von Funktionen einer Variablen durch trigonometrische Reihen mit mehreren einer Variablen proportionalen Argumenten*, Thesis, Dorpat, 1893.
4. H. Bohr, *Zur Theorie der fastperiodischen Funktionen*. I, *Acta Mathematica*, vol. 45 (1925), pp. 29-127. II, *Ibidem*, vol. 46 (1925), pp. 101-214.
5. C. Carathéodory, *Vorlesungen über reelle Funktionen*. 2d edition. Leipzig and Berlin, 1927.
6. E. Cartan, *Sur la structure des groupes de transformations finis et continus*, Thesis, Paris, 1894.
7. E. Cartan, *Les groupes projectifs qui ne laissent invariante aucune multiplicité plane*, *Bulletin de la Société Mathématique de France*, vol. 41 (1913), pp. 53-96.
8. P. J. Daniell, *Differentiation with respect to a function of bounded variation*, *Transactions of the American Mathematical Society*, vol. 19 (1918), pp. 353-362.
9. M. Fréchet, *Pri la funkcio ekvacio* $f(x+y)=f(x)+f(y)$, *L'Enseignement Mathématique*, vol. 15 (1913), pp. 390-393.
10. A. Haar, *Über unendliche kommutative Gruppen*, *Mathematische Zeitschrift*, vol. 33 (1931), pp. 129-159.
11. A. Haar, *Der Massbegriff in der Theorie der kontinuierlichen Gruppen*, *Annals of Mathematics*, (2), vol. 34 (1933), pp. 147-169.
12. G. Hamel, *Eine Basis aller Zahlen und die unstetigen Lösungen der Funktionalgleichung* $f(x+y)=f(x)+f(y)$, *Mathematische Annalen*, vol. 60 (1905), pp. 459-462.
13. F. Hausdorff, *Mengenlehre*. 2d edition. Berlin and Leipzig, 1927.
14. J. v. Neumann, *Über die analytischen Eigenschaften von Gruppen linearer Transformationen und ihrer Darstellungen*, *Mathematische Zeitschrift*, vol. 30 (1929), pp. 3-42.
15. J. v. Neumann, *Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren*, *Mathematische Annalen*, vol. 102 (1930), pp. 49-131.
16. J. v. Neumann, *Zur Algebra der Funktionaloperationen und Theorie der normalen Operatoren*, *Mathematische Annalen*, vol. 102 (1930), pp. 370-427.
17. J. v. Neumann, *Über Funktionen von Funktionaloperatoren*, *Annals of Mathematics*, (2), vol. 32 (1931), pp. 191-226.
18. J. v. Neumann, *Zur Operatorenmethode in der klassischen Mechanik*, *Annals of Mathematics*, (2), vol. 33 (1932), pp. 587-642.
19. J. v. Neumann, *Zum Haarschen Mass in topologischen Gruppen*, *Compositio Mathematica*, vol. 1 (1934), pp. 106-114.

20. E. Schmidt, *Zur Theorie der linearen und nichtlinearen Integralgleichungen*, Mathematische Annalen, vol. 63 (1907), pp. 433–476.
21. I. Schur, *Neue Begründung der Theorie der Gruppencharaktere*, Sitzungsberichte der Preussischen Akademie, Phys. Math. Kl., 1905, pp. 406–432.
22. I. Schur, *Neue Anwendungen der Integralrechnung auf Probleme der Invariantentheorie*, Sitzungsberichte der Preussischen Akademie, Phys. Math. Kl., 1924, pp. 183–208.
23. W. Sierpinski, *Sur l'équation fonctionnelle $f(x+y)=f(x)+f(y)$* , Fundamenta Mathematicae, vol. 1 (1920), pp. 125–129.
24. M. H. Stone, *Linear Transformations in Hilbert Space*, American Mathematical Society Colloquium Publications, vol. XV, 1932.
25. A. Tichonoff, *Über einen Metrisationssatz von P. Urysohn*, Mathematische Annalen, vol. 95 (1926), pp. 139–142.
26. P. Urysohn, *Über die Mächtigkeit der zusammenhängenden Mengen*, Mathematische Annalen, vol. 94 (1925), pp. 262–308.
27. P. Urysohn, *Zum Metrisationsproblem*, Mathematische Annalen, vol. 94 (1925), pp. 309–315.
28. H. D. Ursell, *Normality of almost periodic functions*, First Note, Journal of the London Mathematical Society, vol. 4 (1929), pp. 123–127. Second Note, Ibidem, vol. 5 (1930), pp. 47–50.
29. B. L. van der Waerden, *Stetigkeitssätze der halbeinfachen Lieschen Gruppen*, Mathematische Zeitschrift, vol. 36 (1933), pp. 780–786.
30. H. Weyl, *Theorie der Darstellung kontinuierlicher halbeinfacher Gruppen durch lineare Transformationen. I*, Mathematische Zeitschrift, vol. 23 (1925), pp. 271–309.
31. H. Weyl, *Integralgleichungen und fastperiodische Funktionen*, Mathematische Annalen, vol. 97 (1927), pp. 338–356.
32. H. Weyl and F. Peter, *Die Vollständigkeit der primitiven Darstellungen einer geschlossenen kontinuierlichen Gruppe*, Mathematische Annalen, vol. 97 (1927), pp. 737–755.
33. N. Wiener and R. E. A. C. Paley, *Characters of Abelian groups*, Proceedings of the National Academy of Sciences, vol. 19 (1933), pp. 253–257.

THE DIRAC EQUATION OF PROJECTIVE RELATIVITY

1. In this note we discuss the extension of the Dirac equation to general relativity. In order to have equations which are automatically invariant with respect to gauge as well as coördinate and spin transformations, we use the method of projective relativity. This theory is physically equivalent to the original general relativity theory of an Einstein gravitational and a Maxwell electromagnetic field.

We are led to a class of equations of the Dirac type. One of these reduces exactly to the Dirac equation of a charged particle in special relativity. This equation is identical with the one given by Schrödinger¹ and therefore equivalent to the one given by Fock.² The others contain extra terms which correspond to physical situations in which the field of a dipole is superposed on the field of the charge. Such extra terms, with the charge 0, have been proposed by Pauli in order to explain the properties of the neutron.³ The class thus contains equations first proposed by Schouten and Van Dantzig,⁴ in which an extra term appeared. One of these equations coincides with an equation proposed by Pauli,⁵ and may be considered as the simplest one in the projective notation, while the one without the extra term is simpler in the affine notation (equation (3.3)).

In the following discussions the weight $1/4$ is assigned to spinors, which are to be interpreted as Dirac wave functions in order that expressions of the type,

$$J^\alpha = \gamma_{A\dot{B}}^\alpha \psi^A \dot{\psi}^{B*}, \quad J^{\alpha\beta} = s_{A\dot{B}}^{\alpha\beta} \psi^A \dot{\psi}^{B*}$$

which may represent the current vector and other invariants of the Dirac theory, shall behave as absolute projective vectors and tensors. We are using the notations of previous notes on spinors in these PROCEEDINGS.⁶

The most general equation satisfying the conditions of linearity and invariance is

$$L^\alpha \psi_{, \alpha} + N\psi = 0 \quad (1.1)$$

where the comma refers to projective differentiation, and L^α and N are arbitrary spinors. The projective differentiation is with respect to the projective connection whose components $\Gamma_{\beta\gamma}^\alpha$ are the Christoffel symbols of the second kind formed from the fundamental projective tensor⁷ $\gamma_{\alpha\beta}$. Since any other projective connection associated with $\gamma_{\alpha\beta}$ would differ from this one by a projective tensor, we do not lose any generality by this restriction. The additional terms which would arise from using another projective connection are accounted for by the arbitrary matrix N .

Since the coefficients of the only terms involving $\frac{\partial \psi}{\partial x^i}$ are L^i , these four matrices must reduce to the Dirac matrices in the special relativity case. This is accomplished by identifying them with matrices γ^i which satisfy

$$1/2(\gamma^i \gamma^j + \gamma^j \gamma^i) = 1/4 g^{ij} \cdot 1 \quad (1.2)$$

where $g^{ij} = \gamma^{ij}$ are the contravariant components of the fundamental metric tensor.⁷ In order that $\gamma_{\alpha\beta}$ shall be the fundamental projective tensor of P. R. we choose the matrices γ^α and γ according to the footnote⁷ below rather than according to S. P. R. Equation (1.1) becomes

$$\gamma^i \psi_{, i} + L^0 \psi_{, 0} + N\psi = 0. \quad (1.3)$$

Using the arbitrariness of N , this may be written as

$$\gamma^\alpha \psi_{, \alpha} = \gamma_0 M \psi \quad (1.4)$$

where M is a spinor which we will determine by requiring that equation (1.4) reduce exactly to the usual Dirac equation in the special relativity case.

2. In view of the formula for the projective derivative of a spinor (1.4) is equivalent to

$$\gamma^\alpha \left(\frac{\partial \psi}{\partial x^\alpha} + K_\alpha \psi \right) = \gamma_0 M \psi \quad (2.1)$$

where (P. D. S., p. 88, equation (4.4))

$$K_\alpha = \Gamma_{\nu\alpha}^\mu \gamma_\mu \gamma^\nu + \gamma_\mu \frac{\partial \gamma^\mu}{\partial x^\alpha} + \frac{1}{2} \eta_\alpha \text{ and } \eta_{B\alpha}^A = \gamma^{EA} \frac{\partial \gamma_{EB}}{\partial x^\alpha}. \quad (2.2)$$

From P. R., page 44, we have

$$\begin{aligned} \Gamma_{\nu 0}^\mu &= \gamma^{i\mu} \varphi_{i\nu} & \Gamma_{00}^\alpha &= 0 \\ \Gamma_{jk}^i &= \{j_k^i\} + \varphi_j^i \varphi_k + \varphi_k^i \varphi_j \\ \Gamma_{jk}^0 &= -\varphi_i \Gamma_{jk}^i + \frac{1}{2} \left(\frac{\partial \varphi_j}{\partial x^k} + \frac{\partial \varphi_k}{\partial x^j} \right) \end{aligned} \quad (2.3)$$

where $\{j_k^i\}$ are the Christoffel symbols of the second kind formed from the g_{ij} and $\varphi_{\alpha\beta} = \frac{1}{2} \left(\frac{\partial \varphi_\alpha}{\partial x^\beta} - \frac{\partial \varphi_\beta}{\partial x^\alpha} \right)$.

Substituting (2.3) in (2.2) we obtain

$$K_0 = \varphi_{ij} s^{ij} \quad (2.4)$$

where $s^{\alpha\beta} = \frac{1}{2}(\gamma^\alpha \gamma^\beta - \gamma^\beta \gamma^\alpha)$

$$\gamma^i K_i = -2\gamma_0 K_0 + \gamma^i \varphi_i K_0 + \gamma^i S_i \quad (2.5)$$

where

$$S_i = (\gamma_k - \varphi_k \gamma_0) \left(\{i_k^l\} \gamma^l + \frac{\partial \gamma^k}{\partial x^i} \right) + \frac{1}{2} \eta_i + \gamma_0 \frac{\partial \gamma_0}{\partial x^i}. \quad (2.6)$$

Equation (2.1) becomes

$$\gamma^i \left(\frac{\partial \psi}{\partial x^i} + S_i \psi - I \varphi_i \psi \right) + \gamma_0 (I - K_0) \psi = \gamma_0 M \psi \quad (2.7)$$

where I is the index of ψ and we have made use of the relation

$$\gamma_0 = \gamma_{0\alpha} \gamma^\alpha = \gamma^0 + \varphi_i \gamma^i. \quad (2.8)$$

The spinor M may be expanded in terms of our basis of spinor matrices thus:

$$M = a \cdot 1 + a_\alpha \gamma^\alpha + a_{\alpha\beta} s^{\alpha\beta} \quad (2.9)$$

where a is a projective scalar, a_α a projective vector and $a_{\alpha\beta}$ an anti-symmetric projective tensor.

The quantities a , a_α , $a_{\alpha\beta}$ must be built up out of the physical quantities which belong to the theory, taking account of the prescribed transformation properties.

3. In the special relativity case a coördinate system and spin-frame may be found⁸ in which $\{j_k^i\} = 0$, $\eta_i = 0$, and γ_0 and γ^i are constant

matrices; hence in this coördinate and spin-frame $S_i = 0$. The Dirac equation in this coördinate system and spin-frame is

$$\gamma^i \left(\frac{\partial \psi}{\partial x_i} - \frac{2\pi i}{h} \frac{e}{c} \sqrt{\frac{2}{\kappa}} \varphi_i \psi \right) - \frac{\pi m c}{h} \psi = 0. \quad (3.1)$$

Hence it is clear that (2.7), in which I , a , a_α and $\dot{a}_{\alpha\beta}$ are arbitrary, represents a class of equations of the Dirac type. If we choose

$$a_{\alpha\beta} = b\varphi_{\alpha\beta}, \quad a = I, \quad \text{and} \quad a_\alpha = 4\mu\varphi_\alpha \quad (3.2)$$

equation (2.7) becomes

$$\gamma^i \left(\frac{\partial \psi}{\partial x^i} + S_i \psi - I \varphi_i \psi \right) = \mu \psi + \gamma_0 (b + 1) \varphi_{ij} s^{ij} \psi \quad (3.3)$$

which reduces to the Dirac equation of special relativity if

$$I = \frac{2\pi i}{h} \frac{e}{c} \sqrt{\frac{2}{\kappa}}, \quad \mu = \frac{\pi m c}{h}, \quad \text{and} \quad b = -1. \quad (3.4)$$

Equation (3.3) agrees with the general relativity form of the Dirac equation given by Schrödinger,¹ in all spin-frames in which γ_{AB} and γ_B^A are constants. Our S_i is the traceless part of Schrödinger's Γ_i .

If $b \neq -1$ we obtain an equation with the additional term $\gamma_0 K_0 = \gamma_0 \varphi_{ij} s^{ij}$ multiplied by $b + 1$, which may be interpreted physically as the interaction of the electron with the field of a dipole. Such a term has been suggested by Pauli³ as a possible explanation of the neutron. It also appears in the equation given by Pauli⁵ as the general relativity generalization of the Dirac equation. However, in this work it is multiplied by a factor containing $\sqrt{\kappa}$, the Einstein gravitational constant which makes it negligible. To obtain Pauli's result we set $b = 0$.

4. The affine form of the Dirac equation of general relativity is (3.3) with the choice of constants given in (3.4). The same equation in projective notation is

$$\gamma^\alpha \psi_{, \alpha} = \gamma_0 (I + 4\mu\varphi_\alpha \gamma^\alpha - \varphi_{\alpha\beta} s^{\alpha\beta}) \psi. \quad (4.1)$$

This may be written as

$$\gamma^\alpha (\psi_{, \alpha} - \varphi_\alpha \psi_{, 0}) + 2\gamma_0 K_0 = \mu \psi. \quad (4.2)$$

The equation corresponding to $b = 0$ is

$$\gamma^\alpha (\psi_{, \alpha} - \varphi_\alpha \psi_{, 0}) = \mu \psi. \quad (4.3)$$

This is a very natural one since it is a relation among the terms of the left side of the equation of spin displacement,

$$\psi_{, \alpha} - \varphi_\alpha \psi_{, 0} = 0 \quad (4.4)$$

which corresponds to the projective displacement

$$X_{,\alpha}^{\sigma} - \varphi_{\alpha} X_{,0}^{\sigma} = 0 \quad (4.5)$$

which is used to define the path of a charged particle in projective relativity (P. R., p. 46).

¹ E. Schrödinger, "Diracsches Electron im Schwerfeld 1," *Berl. Ber.*, **11**, 105 (1932).

² V. Fock, *Zeit. Phys.*, **57**, 261 (1929).

³ Cf. J. F. Carlson and J. R. Oppenheimer, *Phys. Rev.*, **41**, 763, 778 (1932).

⁴ Schouten and Van Dantzig, *Ann. Math.*, **34**, 304 (1933).

⁵ W. Pauli, "Die Diracschen Gleichungen für die Materiewellen," *Ann. Phys.*, **18**, 337 (1933).

⁶ We use capital letters, A, B, C ($= 1, 2, 3, 4$) for spin indices, small Roman letters i, j, k ($= 1, 2, 3, 4$) for affine tensor indices and small Greek letters, $\alpha, \beta, \gamma \dots$ ($= 0, 1, 2, 3, 4$) for projective tensor indices. Formulas and results of the following previous papers in these PROCEEDINGS are used: O. Veblen, "Spinors in Projective Relativity," **19**, 979-989 (1933) (referred to as S. P. R.); O. Veblen and A. H. Taub, "Projective Differentiation of Spinors," **20**, 85-92 (1934) (referred to as P. D. S.); J. W. Givens, "Projective Differentiation of Spinors," **20**, 232 (1934).

⁷ For a development of projective relativity see: O. Veblen and B. Hoffman, "Projective Relativity," *Phys. Rev.*, **36**, 810-822 (1930); O. Veblen, *Projektive Relativitätstheorie*, Berlin (1933) (referred to as P. R.).

The fundamental projective tensor of the projective relativity is related to the Einstein gravitational tensor g_{ij} and the electromagnetic projective vector φ_{α} , by the equations

$$\gamma_{ij} = g_{ij} + \varphi_i \varphi_j; \quad \gamma_{0\alpha} = \varphi_{\alpha}; \quad \varphi_0 = 1 \quad (1)$$

of which $g_{ij} = \gamma_{ij}$ is a consequence. The field equations of projective relativity in affine form are (P. R., p. 52)

$$\varphi_{i;s}^s = 0, \quad R^{ij} - \frac{1}{2} g^{ij} R = -2S^{ij} \quad (2)$$

where the semicolon refers to covariant differentiation with respect to g_{ij} , and

$$\varphi_{ij} = \frac{1}{2} \left(\frac{\partial \varphi_i}{\partial x^j} - \frac{\partial \varphi_j}{\partial x^i} \right) \text{ and } S^{ij} = g^{st} \varphi_s^i \varphi_t^j + \frac{1}{4} g^{ij} \varphi_s^s \varphi_t^t. \quad (3)$$

The first set of field equations are the Maxwell equations for free space. They and the remaining field equations agree with those of the general relativity if we set

$$\varphi_i = \sqrt{\frac{\kappa}{2}} \chi_i \quad (4)$$

where κ is the Einstein gravitational constant and χ_i , the electromagnetic four-potential.

The signature of $\gamma_{\alpha\beta}$ must be $(++++)$ or $(----)$ in order to give the right signature for g_{ij} . As Prof. Robertson has pointed out in a lecture which elucidated this whole question, equations (2) are such that changing the sign of the g_{ij} 's would require changing the sign of κ , and the signature $(++++)$ for $\gamma_{\alpha\beta}$ is the one for which κ is positive.

In P. R. the projective tensor satisfies the invariant condition $\gamma_{00} = 1$. This implies that the non-holonomic transformation of period two

$$d\bar{x}^0 = dx^0 + \varphi_i dx^i, \quad d\bar{x}^i = -dx^i \quad (5)$$

will carry it into a form in which

$$\bar{\gamma}_{ij} = g_{ij}, \quad \bar{\gamma}_{i0} = 0, \quad \bar{\gamma}_{00} = 1.$$

The choice of signature $(++++-)$ described in the above paragraph then implies that for any particular point of space-time the coördinates may be chosen so that $g_{ij} = 0$ for $i \neq j$ and $g_{11} = g_{22} = g_{33} = -g_{44} = 1$.

On the other hand in S. P. R. the matrices γ_α were determined in a fixed spin-frame for a projective tensor $\gamma_{\alpha\beta}$ which for the above determination of gauge and coördinates is such that $\gamma_{\alpha\beta} = 0$ if $\alpha \neq \beta$ and $\gamma_{00} = \gamma_{11} = \gamma_{22} = \gamma_{33} = -\gamma_{44} = -1$. A set of matrices γ^α which in the given coördinate gauge and spin-frame give the proper signature for $\gamma_{\alpha\beta}$ (i.e., the signature used in P. R.) are:

$$\begin{aligned} \gamma^0 &= \frac{1}{2} \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix} & \gamma^1 &= \frac{1}{2} \begin{pmatrix} 0 & 0 & i & 0 \\ 0 & 0 & 0 & -i \\ -i & 0 & 0 & 0 \\ 0 & i & 0 & 0 \end{pmatrix} \\ \gamma^2 &= \frac{1}{2} \begin{pmatrix} 0 & 0 & 0 & -1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ -1 & 0 & 0 & 0 \end{pmatrix} & \gamma^3 &= \frac{1}{2} \begin{pmatrix} 0 & 0 & 0 & -i \\ 0 & 0 & -i & 0 \\ 0 & i & 0 & 0 \\ i & 0 & 0 & 0 \end{pmatrix} \\ \gamma^4 &= \frac{1}{2} \begin{pmatrix} -i & 0 & 0 & 0 \\ 0 & -i & 0 & 0 \\ 0 & 0 & i & 0 \\ 0 & 0 & 0 & i \end{pmatrix} \end{aligned}$$

The spin-frame is determined as that in which γ_{AB} and γ_B^A have the form

$$\begin{pmatrix} 0 & i & 0 & 0 \\ -i & 0 & 0 & 0 \\ 0 & 0 & 0 & i \\ 0 & 0 & -i & 0 \end{pmatrix} \quad \text{and} \quad \begin{pmatrix} 0 & -1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & -1 & 0 \end{pmatrix},$$

respectively.

This choice of fundamental matrices changes formula (4.9) of S. P. R. into

$$\gamma_{AB}^* = -\gamma_{BA}^*.$$

* In the special relativity case there exists a coördinate system x in which $g_{ij} = 0$ if $i \neq j$ and $g_{11} = g_{22} = g_{33} = -g_{44} = 1$. To such a coördinate system we apply the non-holonomic transformation (5) and obtain a coördinate system N in which the $\gamma_{\alpha\beta} = 0$ if $\alpha \neq \beta$ and $\gamma_{00} = \gamma_{11} = \gamma_{22} = \gamma_{33} = -\gamma_{44} = 1$. Then in the coördinate system N the matrices γ^α can be made to have the form given above by a suitable spin transformation. We call these constants $\bar{\gamma}^\alpha$. On transforming back to the coördinate system x we obtain

$$\gamma^i = -\bar{\gamma}^i, \quad \gamma_0 = \bar{\gamma}_0 \quad \text{and} \quad \gamma_k = \varphi_k \gamma_0 - \bar{\gamma}_k. \quad (6)$$

Hence with this choice of spin-frame in the special relativity γ^i and γ_0 are constants. Also in the non-holonomic system γ_{AB} is a constant, and since it is invariant under coördinate transformation it is constant in x ; hence $\eta_i = 0$ in x .

ON COMPLETE TOPOLOGICAL SPACES

INTRODUCTION

1. The notion of "completeness" is usually defined only for metric spaces (cf. for instance [1], p. 103). This seems reasonable, because this notion necessarily involves a certain "uniformity of the topology" of the space under consideration. Indeed, the definition of "completeness" is as follows:

DEFINITION I. *If M is a space in which there is defined a metric $\text{dist}(f, g)$ satisfying the usual postulates for distance ([1], p. 94), then a sequence $F: f_1, f_2, \dots$ is fundamental if, for every $\delta > 0$, there exists an $n_1 = n_1(\delta)$ such that $m, n \geq n_1$ imply $\text{dist}(f_m, f_n) < \delta$; and F is convergent if there exists an f such that, for every $\delta > 0$, there exists an $n_2 = n_2(\delta)$ such that $n \geq n_2$ implies $\text{dist}(f, f_n) < \delta$. M is complete if every fundamental sequence is convergent.*

The need of uniformity in M arises from the fact that the elements of a fundamental sequence are postulated to be "near to each other," and not near to any fixed point. As a general topological space (cf. for instance [1], pp. 226–232) has no property which lends itself to the definition of such a "uniformity," it is improbable that a reasonable notion of "completeness" could be defined in it.

However, linear spaces (cf. [1], pp. 95–97, and Definition 1 in this paper), even if only topological, afford a possibility of "uniformization" for their topology: because of their homogeneity everything can be discussed in the neighborhood of 0. Thus one might introduce

DEFINITION I'. *If L is a linear space (cf. above) with a topology (cf. [1], pp. 226–232; of course the "linear" operations of $[\alpha \text{ a real number}]$ and $f+g$ are supposed to be continuous), then a sequence $f_1, f_2, \dots$ is fundamental if, for every neighborhood U of 0 (zero), there exists an $n_1 = n_1(U)$ such that $m, n \geq n_1$ imply $f_m - f_n \in U$;† and convergent if an f can be found so that for every neighborhood U of 0 there exists an $n_2 = n_2(U)$ such that $n \geq n_2$ implies $f - f_n \in U$. L is complete if every fundamental sequence is convergent.*

† The fact that an element x belongs to a set S will be denoted by $x \in S$ (not by $x \subset S$), while $T \subset S$ will mean that the set T is a subset of the set S . Other set-theoretical notations will be used: the sum of a set $(S, T, \dots)$ of sets is $\mathfrak{S}(S, T, \dots)$, the product (that is, the common part of the elements) of $(S, T, \dots)$ is $\mathfrak{P}(S, T, \dots)$, and the complementary set to S is $\mathfrak{C}S$. (These are not the notations of [1].)

This coincides with the previous definition if one defines, as usual, the neighborhoods of a point f_0 in the linear-metric case (in which $\text{dist}(f, g) \equiv \text{dist}(f - g, 0) \equiv \text{dist}(f - g)$) as spheres $S(f_0; \delta): \text{dist}(f, f_0) = \text{dist}(f - f_0) < \delta, \delta > 0$.

Another important notion is "total boundedness" (cf. [1], p. 108). We give the usual definition in the metric case and the generalization for the linear-topological case:

DEFINITION II. *If M is a space in which a metric $\text{dist}(f, g)$ is defined (cf. above), a set $S \subset M$ is totally bounded if, for every $\delta > 0$, there exist a finite number of spheres $S(f_1; \delta), \dots, S(f_n; \delta)$ (of course $n, f_1, \dots, f_n$ all depend on δ) such that $M \subset \bigcup (S(f_1; \delta), \dots, S(f_n; \delta))$.*

DEFINITION II'. *If L is a linear space with a topology (cf. above), a set $S \subset L$ is totally bounded if, for every neighborhood U of 0, a finite number of points $f_1, \dots, f_n$ exist ($n, f_1, \dots, f_n$ all depend on U) such that $M \subset \bigcup (f_1 + U, \dots, f_n + U)$.**

Finally, we repeat the well known definition of "compactness" in a form which is particularly suited for our purposes.

DEFINITION III. *If N is any topological space (cf. above) a set $S \subset N$ is compact if every infinite set $T \subset S$ has a condensation point† $f \in S$. If we require only $f \in N$, this expresses (at least if the countability axiom is satisfied) that S has a compact closure.*

An important fact connecting these notions is that I and II imply III (with $N = M$), the proof resulting from a simple application of the diagonal principle (cf. [1], pp. 108–109). The proof can be transferred immediately to the non-metric case: I' and II' imply III (with $N = L$), provided that the topology of L fulfills Hausdorff's first countability axiom (cf. [1], p. 229, axiom (9); for the proof, cf. Theorem 15 in this paper). That is: if L is complete and fulfills Hausdorff's first countability axiom, every totally bounded set $S \subset L$ has a compact closure.

2. It seems desirable, for various reasons, to get rid of the restriction represented by the countability axiom. Some important examples of linear spaces do not fulfill it.‡ Furthermore, the notions of total boundedness and closure-

* If L is a linear space, we use the following notation: $(f, g \in L; S, T \subset L; \alpha, \beta$ real numbers): αS is the set of all $\alpha f, f \in S$; $f \pm S$ is the set of all $f \pm g, g \in S$; $S \pm T$ is the set of all $f \pm g, f \in S$ and $g \in T$. Note that $\alpha(S+T) = \alpha S + \alpha T$, $(\alpha\beta)S = \alpha(\beta S)$, and $S+T = T+S$, $(S+T)+R = S+(T+R)$; but only the weakened conditions $\alpha S + \beta S \supset (\alpha + \beta)S$, $(S \pm T) \mp T \supset S$ are valid.

† f is a condensation point of T if $\mathfrak{P}(T, U)$ is infinite for every neighborhood U of f .

‡ For example, Hilbert space in its "weak" topology (cf. for instance [2], p. 379); the space of all bounded operators in Hilbert space, in its "strong" and in its "weak" topology (cf. [2], pp. 381–382; for the discussion of all these topologies, [2], pp. 378–388).

compactness play an important role in the general theory of almost periodic functions (cf. the following paper of S. Bochner and J. von Neumann on this subject), and their equivalence is necessary for the smooth working of this theory, which makes no other use of the countability axiom, and which therefore should be workable without its help. (This has actually been done, loc. cit., by the use of the results of this paper.) But if we use the definition of completeness given in I', then II' does not necessarily imply III (that is, total boundedness does not imply closure-compactness) if the countability axiom does not hold. Therefore we have to find another definition of completeness which leads to the desired implication.

The simplest thing is to postulate this directly:

DEFINITION IV. *If L is a linear space with a topology, it is topologically complete if every totally bounded set $S \subset L$ has a compact closure.*

For metric linear spaces, and even for every linear space satisfying the countability axiom, this is equivalent to the usual definition of completeness (I or I'; cf. Theorem 15). The various spaces mentioned at the beginning of this chapter are topologically complete (cf. Theorem 23). The most important property of this notion is, however, that if L is topologically complete, the linear space formed by the functions with a given domain D and with a range $\subset L$ is (if subjected to certain restrictions, like boundedness, etc., cf. Definition 11) topologically complete too. This is rather obvious for the Definitions I and I', but not at all for IV; we will prove it in Theorem 18. All these properties make our notion of topological completeness just as useful for various applications (for instance in the generalized theory of almost periodic functions, as mentioned above), as the usual notion of (metric) completeness, while its range of generality is essentially wider.

We now pass on to the exact exposition of the subject.

I. DEFINITIONS

3. We define linear spaces in the usual way (cf. for instance [1], pp. 95-97):

DEFINITION 1. *The set L is a linear space, if, for $f, g \in L$ and any real number α , αf and $f+g$ are in L and are defined so that*

- | | |
|---|--|
| (1) $f + g = g + f,$ | (2) $(f + g) + h = f + (g + h),$ |
| (3) $1 \cdot f = f,$ | (4) $\alpha(\beta f) = (\alpha\beta)f,$ |
| (5) $(\alpha + \beta)f = \alpha f + \beta f,$ | (6) $\alpha(f + g) = \alpha f + \alpha g,$ |
| (7) $f + h = g + h$ implies $f = g.$ | |

* It would be sufficient to admit only rational α 's.

By (3) and (5), $f+0 \cdot f=f$; by (1) and (2), $(f+g)+0 \cdot f=f+g$; by interchanging f and g , $(f+g)+0 \cdot g=f+g$; by (7), $0 \cdot f=0 \cdot g$. Thus $0 \cdot f$ is independent of f ; and we call it 0 (zero). We have $f+0=f$. Writing $-f$ for $(-1) \cdot f$, (5) gives $f+(-f)=0$; writing $f-g$ for $f+(-g)$, (1) and (2) give $(f-g)+g=f$; and by (7), $x=f-g$ is the only solution of $x+g=f$. Now all rules of computation for 0 , $-f$, $f-g$ are easily deduced.

A metric (or an "absolute value") in L is defined in the usual way (cf. [1], p. 97):

DEFINITION 2a. *The linear set L is metric if, for every $f \in L$, a real number $\|f\|$, its "absolute value," is defined, such that*

$$(1) \quad \|f\| > 0 \text{ if } f \neq 0, \quad (2) \quad \|\alpha f\| = |\alpha| \cdot \|f\|, \quad (3) \quad \|f+g\| \leq \|f\| + \|g\|.$$

The metric is then defined by $\text{dist}(f, g) = \|f-g\|$.

This $\text{dist}(f, g)$ possesses the characteristic properties of a distance and can be used to define a topology in L (cf. [1], p. 94; also the end of paragraph 1 of this paper). However, we shall not assume that L is metric, but only that it has a topology. This is done in the following definition, in which it was attempted to reduce the strength of the postulates to the necessary minimum.

DEFINITION 2b. *The linear set L is topological if a set $\mathfrak{U}$ of sets $U \subset L$ is given such that*

- (1) *if $U \in \mathfrak{U}$, then $0 \in U$,*
- (2) *there is a sequence $U_1, U_2, \dots \in \mathfrak{U}$ such that $\mathfrak{P}(U_1, U_2, \dots) = (0), \dagger\dagger$*
- (3) *if $U, V \in \mathfrak{U}$, there is a $W \in \mathfrak{U}$ with $W \subset \mathfrak{P}(U, V), \ddagger$*
- (4) *if $U \in \mathfrak{U}$, there is a $V \in \mathfrak{U}$ such that, for every α with $-1 \leq \alpha \leq 1$, $\alpha V \subset U, \S$*
- (5) *if $U \in \mathfrak{U}$, there is a $V \in \mathfrak{U}$ with $V+V \subset U, \S$*
- (6) *if $f \in L$, $U \in \mathfrak{U}$, there is an α with $f \in \alpha U. \S$*

L is "convex" if the further condition

- (7) *if $U \in \mathfrak{U}$, then $U+U \subset 2U \S$*

is fulfilled.

$\dagger$ If Hausdorff's first countability axiom holds, (2) is fulfilled by choosing a complete system of neighborhoods of 0 for $U_1, U_2, \dots$ (cf. [1], p. 229, axiom (9)), but the converse is not true: (2) is essentially weaker than the countability axiom. This is shown by the examples of Part IV, Theorem 23; cf. [3], p. 264.

$\ddagger$ (2) and (3) could be replaced by two other postulates (2') and (3') which extend (2) and restrict (3):

(2') there is an aleph $\aleph^*$ and a set $(U, V, \dots) \subset \mathfrak{U}$ with this aleph $\aleph^*$, such that $\mathfrak{P}(U, V, \dots) = (0)$;

(3') for each set $(U, V, \dots) \subset \mathfrak{U}$ with an aleph $< \aleph^*$ there is a $W \in \mathfrak{U}$ with $W \subset \mathfrak{P}(U, V, \dots)$.

All our discussions could be carried through, with little change, on this basis. In the present form of (2) and (3), $\aleph^* = \aleph_0$ (the set $(U, V, \dots)$ is countable).

$\S$ See first footnote on p. 509

If L is metric, we consider the spheres ($\delta > 0$)

$S^0(f_0; \delta)$: the set of all f with $\|f - f_0\| < \delta$,

$S^1(f_0; \delta)$: the set of all f with $\|f - f_0\| \leq \delta$.

Choosing $\mathfrak{U}$ as the set of all $S^0(0; \delta)$ or as the set of all $S^1(0; \delta)$ makes L topological and convex (one easily verifies Definition 2b, (1)–(7)), and the topology, which we will define with the aid of $\mathfrak{U}$ in Definition 4 and Theorem 6, coincides in this case with the usual metric topology of L .

II. GENERAL THEOREMS

4. In the following discussion of Chapters II and III, L is assumed merely to be a topological space, that is, to fulfill Definition 2b, (1)–(6), except where the contrary is expressly stated.

THEOREM 1. *If $U \in \mathfrak{U}$, $A > 0$, $n = 1, 2, \dots$, there is, for each value of n , a $V \in \mathfrak{U}$ such that, for all sets $\alpha_1, \dots, \alpha_n$ with $-A \leq \alpha_1, \dots, \alpha_n \leq A$, $\alpha_1 V + \dots + \alpha_n V \subset U$.*

Increasing A and n strengthens the statement, so we may assume $A = 2^p$, $n = 2^q$, $p, q = 0, 1, 2, \dots$. It is sufficient to obtain $AV + \dots + AV \subset U$ (n addends), because the W of Definition 2b, 4 (with $\alpha W \subset V$ for $-1 \leq \alpha \leq 1$) will then have the desired properties. As $AV \subset V + \dots + V$ (A addends),* the statement is strengthened if we replace A, n by $1, An = 2^{p+q}$, respectively. Thus we may assume $A = 1$, $n = 2^r$, $r = 0, 1, 2, \dots$. We have to find a $V \in \mathfrak{U}$ with $V + \dots + V \subset U$ (2^r addends).

For $r = 0$ we may choose $V = U$; if we have a V for any $r = 0, 1, 2, \dots$, we can apply Definition 2b, (5), to it, and thus obtain a V for $r + 1$. This completes the proof.

DEFINITION 3. *If $S \subset L$, S_i is the set of all f for which a $U \in \mathfrak{U}$ with $f + U \subset S$ exists.†*

THEOREM 2. $S_{ii} = S_i \subset S$.

$0 \in U$, $f \in f + U$, so that $S_i \subset S$. Therefore $S_{ii} \subset S_i$. Now if $f \in S_i$, that is, if $f + U \subset S$, choose a $V \in \mathfrak{U}$ with $V + V \subset U$ (Definition 2b, (5)). Then for $g \in f + V$, $g + V \subset f + V + V \subset f + U \subset S$, and $g \in S_i$. Thus $f + V \subset S_i$, $f \in S_{ii}$, and therefore $S_i \subset S_{ii}$. This completes the proof.

THEOREM 3. $0 \in U_i$; $(f + S)_i = f + S_i$; if $\alpha \neq 0$, $(\alpha S)_i = \alpha S_i$; $S_i + T_i \subset (S + T)_i$.

The first two statements are obvious. If $f \in S_i$ and $f + U \subset S$, then $\alpha f + \alpha U \subset \alpha S$, and if V is chosen by Theorem 1 with $V/\alpha \subset U$, $\alpha f + V \subset \alpha S$ and

* See first footnote on p. 509

† S_i is the set of inner points of S (cf. Theorem 4).

$\alpha f \in (\alpha S)_i$. Thus $\alpha S_i \subset (\alpha S)_i$; substitution of $1/\alpha$, αS for α , S leads to the result that $(\alpha S)_i \subset \alpha S_i$, proving the third statement. If $f \in S_i$ and $g \in T_i$, $f + g \in f + T_i = (f + T)_i \subset (S + T)_i$, proving the fourth statement.

DEFINITION 4. S is open if $S = S_i$; S is closed if $\mathfrak{CS}^*$ is open. This means that if S is closed, $S = S_{cl}$, where we define $S_{cl} = \mathfrak{C}((\mathfrak{CS})_i)$.†

THEOREM 4. S_i is open and the greatest open set $\subset S$, S_{cl} is closed and the smallest closed set $\supset S$.

The statements about S_i follow from Theorems 2 and 3; those about S_{cl} result from considering $\mathfrak{CS}$.

THEOREM 5. For every S , $S_{cl} = \mathfrak{P}(S + U)$, where U runs over all elements of $\mathfrak{U}$.

If $U \in \mathfrak{U}$, there is a $V \in \mathfrak{U}$ with $-V \subset U$ (Definition 2b, (4)), and thus $V \subset -U$. For this reason $\mathfrak{P}(S + U) = \mathfrak{P}(S - V)$ (U, V run over all $\mathfrak{U}$). Now $f \in \mathfrak{P}(S - V)$ means that for every $V \in \mathfrak{U}$, $f \in g - V$ for some $g \in S$, that is, $g \in S$, $g \in f + V$. But this is equivalent to $f \in S_{cl}$.

THEOREM 6. In the sense in which Hausdorff defined a topology, using the open sets as fundamental notions (cf. [1], p. 228), Definition 4 describes a regular Hausdorff topology, that is, one which fulfills Hausdorff's axioms (1)–(3) and (6) (cf. [1], pp. 228–229; thus all axioms (1)–(6) are fulfilled).

Ad (1): 0 and L are obviously open. Ad (2): If S_1, S_2 are open, assume $f \in \mathfrak{P}(S_1, S_2)$. Then $f + U_1 \subset S_1$, $f + U_2 \subset S_2$, and so with $V \in \mathfrak{U}$, $V \subset \mathfrak{P}(U_1, U_2)$ (by Definition 2b, (3)) and $f + V \subset \mathfrak{P}(S_1, S_2)$, proving that $f \in \mathfrak{P}(S_1, S_2)_i$. Thus $\mathfrak{P}(S_1, S_2)$ is open. Ad (3): If $S, T, \dots$ are open, $\mathfrak{S}(S, T, \dots)$ is obviously open. Ad (6): Let $\bar{S}$ be closed, f not an element of $\bar{S}$. Then $\mathfrak{CS}$ is open, $f \in \mathfrak{CS}$, so that $U \in \mathfrak{U}$, $f + U \subset \mathfrak{CS}$. Choosing $V \in \mathfrak{U}$ with $V + V \subset U$ (by Definition 2b, (3)) we have $(f + V_i)_{cl} \subset (f + V)_{cl} \subset f + V + V$ (by Theorem 5) $\subset f + U \subset \mathfrak{CS}$, that is, for $T = f + V_i$, T is open, $f \in T$, and $\mathfrak{P}(T_{cl}, \bar{S})$ is empty.

In the following discussions we shall always consider L as topologized by the topology of Definition 4 and Theorem 6, except where the contrary is explicitly stated. Thus we can use the whole topological terminology: we can speak of open and closed sets, which have already been defined in harmony with this by Definition 4, of continuous functions, limits of sequences, condensation points, etc.

THEOREM 7. αf and $f + g$ are continuous functions of α , f and f, g respectively.

* See second footnote on p. 508.

† S_{cl} , the closure of S , is the set of all points and all condensation points of S (cf. Theorem 4).

Ad $f+g$: Assume $f_0+g_0 \in S$, S open. Choose $U \in \mathcal{U}$, $f_0+g_0+U \subset S$, $V \in \mathcal{U}$, $V+V \subset U$. Then

$$(f_0 + V_i) + (g_0 + V_i) \subset f_0 + g_0 + V_i + V_i \subset f_0 + g_0 + V + V \subset f_0 + g_0 + U \subset S,$$

that is, the sets $T_1=f_0+V_i$, $T_2=g_0+V_i$ are open, $f_0 \in T_1$, $g_0 \in T_2$, $T_1+T_2 \subset S$.

Ad αf : Due to the continuity of $f+g$ (and Definition 1, (5)–(6)) we need consider only $\alpha_0=0$ or $f_0=0$. Then, in any event, $\alpha_0 f_0=0$, so that we have the situation $0 \in S$, S open. Ad $\alpha_0=0$: Choose $U \in \mathcal{U}$, $U \subset S$, and $V \in \mathcal{U}$ by Theorem 1 with $n=2$, $A=1$. Now choose a β with $f_0 \in \beta V$ (by Definition 2b, (6)); then for

$$|\alpha| < \frac{1}{\max(1, |\beta|)}$$

we have $\alpha(f_0+V) \subset \alpha\beta V + \alpha V \subset U \subset S$, that is, the sets

$$A: \quad |\alpha| < \frac{1}{\max(1, |\beta|)}$$

and $T=f_0+V_i$ are open, $0 \in A$, $f_0 \in T$, $AT \subset S$. Ad $f_0=0$: Choose $U \in \mathcal{U}$, $U \subset S$ and $V \in \mathcal{U}$ by Theorem 1 with $n=1$, $A=|\alpha_0|+1$. Then $|\alpha-\alpha_0| < 1$ implies $|\alpha| \leq A$, $\alpha V \subset U \subset S$, that is, the sets $A: |\alpha-\alpha_0| < 1$ and $T=V_i$ are open, $\alpha_0 \in A$, $0 \in T$, $AT \subset S$.

5. We now introduce the notion of boundedness, which is usually considered as a metric notion. One sees at once, remembering the remarks made at the end of §3, that our general definition coincides in the case of a metric L with the usual definition.

DEFINITION 5. *S is bounded if, for every $U \in \mathcal{U}$, there is an α such that $S \subset \alpha U$.*

Remark. We could replace herein the $U \in \mathcal{U}$ by the open sets T with $0 \in T$ for if T is such a set, a $U \in \mathcal{U}$, $U \subset T$, exists, and for $U \in \mathcal{U}$ an open T , $0 \in T$, $T \subset U$ exists: $T = U_i$.

THEOREM 8. *Every finite set is bounded. If $S, \dots, T$ are a finite number of bounded sets, $\oplus(S, \dots, T)$ is bounded. If S, T are bounded sets, $\alpha S, f+S, S+T$ are bounded.*

It follows by Definition 2b, (6), that a one-element set $\{f\}$ is bounded; therefore every finite set is bounded if the second statement is true. If the second statement holds for two addends, it holds by induction for any finite number. Let S, T be bounded, $U \in \mathcal{U}$, choose $V \in \mathcal{U}$ by Definition 2b, (4), and α, β with $S \subset \alpha V$, $T \subset \beta V$. Then for $\gamma = \max(|\alpha|, |\beta|)$,

$$S \subset \alpha V = \gamma \left(\frac{\alpha}{\gamma} V \right) \subset \gamma U;$$

similarly, $T \subset \gamma U$, $\mathfrak{S}(S, T) \subset \gamma U$. Thus the first two statements are proved.

In the last statement the first part (concerning αS) is obvious; the second follows from the third by putting $T = (f)$, so that we have to prove only the third part. Assume $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}$ by Theorem 1 with $n=2$, $A=1$. Choose β, γ with $S \subset \beta V$, $T \subset \gamma V$. Then, for $\delta = \max(|\beta|, |\gamma|)$,

$$S + T \subset \beta V + \gamma V = \delta \left(\frac{\beta}{\delta} V \right) + \delta \left(\frac{\gamma}{\delta} V \right) \subset \delta U.$$

This completes the proof.

The next definition is a repetition of Definition II' in §1 (cf. the comment given there).

DEFINITION 6. *S is totally bounded if, for every $U \in \mathfrak{U}$, there is a finite number of elements $f_1, \dots, f_n$ of L , such that $S \subset \mathfrak{S}(f_1 + U, \dots, f_n + U)$.*

Remark 1. The $f_1, \dots, f_n$ could be restricted to S . In this form the condition is obviously sufficient; but it is also necessary: Choose $V \in \mathfrak{U}$ by Theorem 1, with $n=2$, $A=1$; then $V - V \subset U$. Apply Definition 6 to V : $S \subset \mathfrak{S}(g_1 + V, \dots, g_n + V)$. As the sets $g_r + V$ which contain no point of S can be omitted, we may assume that every $g_r + V$ contains a point of S , say f_r . Then $g_r \in f_r - V$, $S \subset \mathfrak{S}(f_1 - V + V, \dots, f_n - V + V) \subset \mathfrak{S}(f_1 + U, \dots, f_n + U)$.

Remark 2. For the reasons given in the Remark after Definition 5, we could replace the $U \in \mathfrak{U}$ by the open sets T with $0 \in T$ in all these considerations.

Remark 3. The set of all real numbers is a particularly simple linear space. It is clear that its customary metric, $\|x\|$ = absolute value of x , is a metric in the sense of Definition 2a, and a topology in the sense of Definition 2b (cf. the end of Part I). The boundedness in the sense of Definition 5, and the total boundedness in the sense of Definition 6, obviously coincide with the customary notion of boundedness for real numbers in this case.

THEOREM 9. *Every finite set is totally bounded. The set of all αf , $-1 \leq \alpha \leq 1$ (f fixed), is totally bounded. If S, T are totally bounded, αS , $f + S$, $S + T$ are also totally bounded, and so is $\mathfrak{S}(S, T)$.*

If $L_1, \dots, L_k, L$ are topological linear spaces, if $S_\kappa \subset L_\kappa$ and is totally bounded, $\kappa=1, \dots, k$, if $\mathfrak{F}(f_1, \dots, f_k)$ is a function with the domain $f_\kappa \in S_\kappa$, $\kappa=1, \dots, k$, uniformly continuous in this domain, and with a range $\subset L$, then the set of all $\mathfrak{F}(f_1, \dots, f_k)$, $f_\kappa \in S_\kappa$, $\kappa=1, \dots, k$, is totally bounded.

The statements for finite sets, $f + S$, $\mathfrak{S}(S, T)$, are obvious, and the statement for αS follows from Theorem 1. The statements concerning the sets αf , $-1 \leq \alpha \leq 1$, and $S + T$ are both special cases of the last statement ($k=1$, L = set of all real numbers, $L_1 = L$, f_1 is to be replaced by α , $\mathfrak{F}(\alpha) = \alpha f$; and $k=2$, $L_1 = L_2 = L$, $\mathfrak{F}(f_1, f_2) = f_1 + f_2$ respectively; cf. Theorem 7 and Remark

3 after Definition 6). So our only task is to prove that statement. Denote the $\mathfrak{U}$ of $L_1, \dots, L_k, L$ by $\mathfrak{U}_1, \dots, \mathfrak{U}_k, \mathfrak{U}$ respectively. If a $U \in \mathfrak{U}$ is given, the uniform continuity of $\mathfrak{F}$ means that we can choose $U_1 \in \mathfrak{U}_1, \dots, U_k \in \mathfrak{U}_k$ so that the conditions $f_\kappa - g_\kappa \in U_\kappa, f_\kappa, g_\kappa \in S_\kappa, \kappa=1, \dots, k$, imply that $\mathfrak{F}(f_1, \dots, f_k) - \mathfrak{F}(g_1, \dots, g_k) \in U$. Now apply Definition 6 and its Remark 1 to S_κ : $S_\kappa \subset \mathfrak{S}(g_{\kappa,1} + U_\kappa, \dots, g_{\kappa,n_\kappa} + U_\kappa), g_{\kappa,1}, \dots, g_{\kappa,n_\kappa} \in S_\kappa$. If a system $f_\kappa \in S_\kappa, \kappa=1, \dots, k$, is given, we have $f_\kappa \in g_{\kappa,\nu_\kappa} + U_\kappa$, with a $\nu_\kappa=1, \dots, n_\kappa$ for each $\kappa=1, \dots, k$, and thus $\mathfrak{F}(f_1, \dots, f_k) \in \mathfrak{F}(g_{1,\nu_1}, \dots, g_{k,\nu_k}) + U$. So if we put $n=n_1 \dots n_k$, and arrange the $\mathfrak{F}(g_{1,\nu_1}, \dots, g_{k,\nu_k})$ ($\nu_\kappa=1, \dots, n_\kappa, \kappa=1, \dots, k$) in some order $\mathfrak{F}_1, \dots, \mathfrak{F}_n$, then our set is contained in $\mathfrak{S}(\mathfrak{F}_1 + U, \dots, \mathfrak{F}_n + U)$.

THEOREM 10. *Every totally bounded set S is bounded.*

Let S be totally bounded, $U \in \mathfrak{U}$. Choose $V \in \mathfrak{U}$ by Theorem 1, with $n=2, A=1$, and then $f_1, \dots, f_n$ with $S \subset \mathfrak{S}(f_1 + V, \dots, f_n + V)$. Select $\alpha_1, \dots, \alpha_n$ with $f_\kappa \in \alpha_\kappa V$ and put $\beta = \max(1, |\alpha_1|, \dots, |\alpha_n|)$; then

$$f_\kappa + V \subset \alpha_\kappa V + V = \beta \left(\frac{\alpha_\kappa}{\beta} V \right) + \beta \left(\frac{1}{\beta} V \right) \subset \beta U.$$

THEOREM 11. *The boundedness (total boundedness) of S is a necessary and sufficient condition for that of S_{cl} .*

As $S \subset S_{\text{cl}}$, the condition is necessary. If $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}, V + V \in \mathfrak{U}$; then by Theorem 5, $V_{\text{cl}} \subset V + V \subset U$, and thus $S \subset \alpha V$ (or $S \subset \mathfrak{S}(f_1 + V, \dots, f_n + V)$) implies $S_{\text{cl}} \subset \alpha U$ (or $S_{\text{cl}} \subset \mathfrak{S}(f_1 + U, \dots, f_n + U)$). Thus the condition is also sufficient. †

We now investigate convexity.

THEOREM 12. *If L is convex (Definition 2b, (7)) then, for $U \in \mathfrak{U}$ and $\alpha_1, \dots, \alpha_n \geq 0, \alpha_1 U_{\text{cl}} + \dots + \alpha_n U_{\text{cl}} = (\alpha_1 + \dots + \alpha_n) U_{\text{cl}}$.*

Induction proves this theorem for all $n=1, 2, \dots$, if it holds for $n=2$. For $\alpha_1 = \alpha_2 = 0$ it is obvious, therefore we may assume $\alpha_1 + \alpha_2 > 0$. Division by $\alpha_1 + \alpha_2$ then gives

$$\alpha U_{\text{cl}} + (1 - \alpha) U_{\text{cl}} = U_{\text{cl}} \quad \left(\alpha = \frac{\alpha_1}{\alpha_1 + \alpha_2}, 0 \leq \alpha \leq 1 \right).$$

As $\supset$ is obvious, we need to prove only $\subset$.

Now Definition 2b, (7), states that $U + U \subset 2U$; iterated n times, this becomes $U + \dots + U$ (2^n addends) $\subset 2^n U$, and, a fortiori, $kU + (2^n - k)U \subset 2^n U$, $k=0, 1, \dots, 2^n$. Thus $\alpha U + (1 - \alpha)U \subset U$ if $0 \leq \alpha \leq 1$ with α dyadic-rational; from this, consideration of continuity leads to $\alpha U_{\text{cl}} + (1 - \alpha)U_{\text{cl}} \subset U_{\text{cl}}$ for all $0 \leq \alpha \leq 1$, completing the proof.

DEFINITION 7. S_{conv} is the set of all $\alpha_1 f_1 + \dots + \alpha_n f_n$, with $n=1, 2, \dots$, $\alpha_1, \dots, \alpha_n \geq 0$, $\alpha_1 + \dots + \alpha_n = 1$, $f_1, \dots, f_n \in S$.

THEOREM 13. If L is convex, then, for all $U \in \mathfrak{U}$, $(U_{\text{ol}})_{\text{conv}} = U_{\text{ol}}$.

This follows immediately from Theorem 12.

THEOREM 14. The boundedness (total boundedness) of S is a necessary condition for that of S_{conv} ; if L is convex, it is also sufficient.

As $S \subset S_{\text{conv}}$, the condition is necessary. For the sufficiency we have to assume that L is convex. If $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}$, $V + V \subset U$; then $V_{\text{conv}} \subset (V_{\text{ol}})_{\text{conv}} = V_{\text{ol}} \subset V + V \subset U$ (by Theorems 5, 13). Thus $S \subset \alpha V$ implies $S_{\text{conv}} \subset \alpha V_{\text{conv}} \subset \alpha U$, settling the case for boundedness. $S \subset \mathfrak{S}(f_1 + V, \dots, f_n + V)$ implies $S_{\text{conv}} \subset \mathfrak{S}(\beta_1 f_1 + \dots + \beta_n f_n + V_{\text{conv}}) \subset \mathfrak{S}(\beta_1 f_1 + \dots + \beta_n f_n + U)$, where the $\beta_1, \dots, \beta_n$ run over all combinations with $\beta_1, \dots, \beta_n \geq 0$, $\beta_1 + \dots + \beta_n = 1$. If the set of these $\beta_1 f_1 + \dots + \beta_n f_n$ is totally bounded, then this set is contained in $\mathfrak{S}(g_1 + U, \dots, g_p + U)$, and thus $S_{\text{conv}} \subset \mathfrak{S}(g_1 + U + U, \dots, g_p + U + U)$. We could replace U by V , and we would then have $S_{\text{conv}} \subset \mathfrak{S}(g_1 + U, \dots, g_p + U)$. This settles the case for total boundedness too, provided that the set of the $\beta_1 f_1 + \dots + \beta_n f_n$ (n and $f_1, \dots, f_n$ fixed) is totally bounded.

This is a subset of the $(\beta_1 f_1 + \dots + \beta_n f_n)$ -set with $0 \leq \beta_1 \leq 1, \dots, 0 \leq \beta_n \leq 1$. This set is disposed of by Theorem 9, if each set $\beta_j f_j$, $0 \leq \beta_j \leq 1$, is totally bounded, but this again follows from Theorem 9.

Finally we repeat Definition III in §1.

DEFINITION 8. S is compact if every infinite set $T \subset S$ has a condensation point $f \in S$.

Note that we did not assume that any of Hausdorff's countability axioms hold (cf. [1], p. 229, axioms (9) and (10)), and that we still are considering "compactness" and not the Alexandroff-Urysohn "bicomcompactness" (cf. [3], [4], in particular [3], pp. 259–260), although the latter is specially adapted to these cases. The reason is that for the totally bounded sets S compactness implies the countability axioms (although they need not hold for L , cf. Theorem 16) and thus bicomcompactness.

III. TOPOLOGICAL COMPLETENESS

6. The two definitions of completeness which we discussed in §1, I', and IV, are the following:

DEFINITION 9. L is sequentially complete if every "fundamental sequence" $f_1, f_2, \dots$ in L (i.e., every sequence such that for each $U \in \mathfrak{U}$ there is an $n_1 = n_1(U)$, such that $m, n \geq n_1$ imply $f_m - f_n \in U$) is "convergent" (i.e., an f exists such that for each $U \in \mathfrak{U}$ there is an $n_2 = n_2(U)$, such that $n \geq n_2$ implies $f_n - f \in U$).

DEFINITION 10. *L is topologically complete if every closed and totally bounded set $S \subset L$ is compact.*

To characterize the relationship between these two notions, we prove

THEOREM 15. *Sequential completeness is a necessary condition for topological completeness; if L satisfies Hausdorff's first countability axiom (cf. above), it is also sufficient.*

Assume first that L is topologically complete, and let $f_1, f_2, \dots$ be a fundamental sequence. If $U \in \mathcal{U}$, $f_m \in f_{N_1} + U$ for $m \geq N_1 = N_1(U)$; thus $(f_1, f_2, \dots) \subset \mathcal{S}(f_1 + U, \dots, f_{N_1} + U)$. So $(f_1, f_2, \dots)$ is totally bounded; $(f_1, f_2, \dots)_{\alpha_1}$ is also totally bounded (by Theorem 11) and closed, and thus compact. If infinitely many $f_1, f_2, \dots$ are different, $(f_1, f_2, \dots)$ is an infinite set $\subset (f_1, f_2, \dots)_{\alpha_1}$, and has a condensation point f , that is, an f such that, for each $U \in \mathcal{U}$, there are infinitely many n for which $f_n \in f + U$. If only a finite number of $f_1, f_2, \dots$ are different, infinitely many f_n must coincide, and their common value f has the above property. So such an f exists in any event. Corresponding to $U \in \mathcal{U}$, choose $V \in \mathcal{U}$, $V + V \subset U$, $n_1 = n_1(V)$, and the above f with respect to V . Then $f_n - f \in V$ occurs for some $n \geq n_1$, and $f_m - f_n \in V$ for every $m, n \geq n_1$; thus, for every $m \geq n_1$, $f_m - f = (f_m - f_n) + (f_n - f) \in V + V \subset U$. Putting $n_2(U) = n_1(V)$ we see that $f_1, f_2, \dots$ is convergent, and thus L is sequentially complete.

Now consider the converse situation. Let L be sequentially complete, let $U_1, U_2, \dots$ be a complete system of neighborhoods for 0 (cf. [1], p. 229; this means that for each $U \in \mathcal{U}$ some $U_n \subset U$), and assume S to be totally bounded and closed. As every infinite $T \subset S$ contains a sequence $f_1, f_2, \dots$ of distinct elements, we may assume $T = (f_1, f_2, \dots)$. As every fundamental sequence is convergent and thus has a condensation point f , we need only exhibit a fundamental subsequence $f^{(1)}, f^{(2)}, \dots$ in a sequence $(f_1, f_2, \dots) \subset S$. Now we can apply the well known diagonal process.

Choose $g_{n1}, \dots, g_{nmn}$ with $S \subset (g_{n1} + U_n, \dots, g_{nmn} + U_n)$. Some $g_{1\nu} + U_1$, say $\nu = \nu_1$, contains infinitely many elements $f_1^{(1)}, f_2^{(1)}, \dots$ of the sequence $f_1, f_2, \dots$. Some $g_{2\nu} + U_2$, say $\nu = \nu_2$, contains infinitely many elements $f_1^{(2)}, f_2^{(2)}, \dots$ of the sequence $f_1^{(1)}, f_2^{(1)}, \dots$. And so on. Now put $f^{(1)} = f_1^{(1)}$, $f^{(2)} = f_2^{(2)}$, $\dots$. If $U \in \mathcal{U}$, choose $V \in \mathcal{U}$, $V - V \subset U$ (cf. Definition 2b, (4), (5)) and $U_p \subset V$. For $m \geq p$, $f^{(m)} = f_m^{(m)}$ belongs to $(f_1^{(m)}, f_2^{(m)}, \dots)$, thus to $(f_1^{(p)}, f_2^{(p)}, \dots)$, and to $g_{mp} + U_p$. Thus, for $m, n \geq p$, $f^{(m)} - f^{(n)} \in U_p - U_p \subset V - V \subset U$. Putting $n_1(U) = p$ we see that $f^{(1)}, f^{(2)}, \dots$ is fundamental, thus completing the proof of the topological completeness of L .

As our main interest belongs to the cases in which L violates the countability axiom, and as the equivalence of sequential and of topological com-

pleteness has not been established for them, we continue by investigating the properties of topological completeness.

THEOREM 16. *If L is topologically complete and S totally bounded, then if we consider S as a space with the topology of L , this is a normal and separable Hausdorff topology in S , that is, one which fulfills Hausdorff's axioms (1)–(3), (8), and (10) (cf. [1], pp. 228–229; thus all axioms (1)–(10) are fulfilled).*

Remark. For this reason we can replace condensation points in S by limits of convergent sequences.—

As $S \subset S_{e1}$ and as S_{e1} is also totally bounded, we may consider S_{e1} instead of S , that is, we can restrict ourselves to closed sets S . Then S is compact. The topology of L fulfilled axioms (1)–(3), (6) in L , therefore it also fulfills them in S . Let us therefore consider the other axioms.

Ad (9): Choose $U_1, U_2, \dots$ by Definition 2b, (2), with $\mathfrak{P}(U_1, U_2, \dots) = (0)$. Now choose (by Definition 2b, (3), (5)) $V_n \in \mathfrak{U}, V_n + V_n \subset U_n$ and $W_n \in \mathfrak{U}, W_n \subset \mathfrak{P}(V_1, \dots, V_n)$. Let T be an open set, $0 \in T$. Assume that, for $n=1, 2, \dots$, $\mathfrak{P}(W_n, S)$ is not a subset of $\mathfrak{P}(T, S)$. Then there exists an f_n with $f_n \in \mathfrak{P}(W_n, S)$, f_n not an element of $\mathfrak{P}(T, S)$. If infinitely many $f_1, f_2, \dots$ are different, $(f_1, f_2, \dots)$ is an infinite set $\subset S$, and has a condensation point f , that is, an f such that, for each $U \in \mathfrak{U}$, there are infinitely many m for which $f_m \in f + U$. If only a finite number of $f_1, f_2, \dots$ are different, infinitely many f_n must coincide, and their common value f has the above property. So such an f exists in any event. Choose $m \geq n$ and $f_m \in f + U$. $f_m \in W_m \subset V_n$, $f_m \in \mathfrak{U}T$, so that f is a point or a condensation point of V_n and of $\mathfrak{U}T$; hence $f \in (V_n)_{e1} \subset V_n + V_n \subset U_n$ (by Theorem 5), and $f \in \mathfrak{U}T$ ($\mathfrak{U}T$ is closed). Thus $f \in \mathfrak{P}(U_1, U_2, \dots) = (0)$, $f = 0 \in T$, contradicting the condition that $f \notin \mathfrak{U}T$. This proves that an $n=1, 2, \dots$ with $\mathfrak{P}(W_n, S) \subset \mathfrak{P}(T, S)$ must exist, so that $(W_1)_i, (W_2)_i, \dots$ form a complete system of neighborhoods of 0 in S , and $f + (W_1)_i, f + (W_2)_i, \dots$ form a complete system of neighborhoods of f in S (consider 0 in $-f + S$).

Ad (10): Take the $W_1, W_2, \dots$ constructed above and choose $X_n \in \mathfrak{U}, X_n - X_n \subset W_n$ (by Theorem 1, $n=2, A=1$, and then $g_{n1}, \dots, g_{nmn}$ with $S \subset \mathfrak{G}(g_{n1} + (X_n)_i, \dots, g_{nmn} + (X_n)_i)$. If an open set T and an $f \in \mathfrak{P}(S, T)$ are given, there is an n for which $\mathfrak{P}(S, f + W_n) \subset \mathfrak{P}(S, T)$, and a ν with $f \in g_{n\nu} + (X_n)_i$. Then $g_{n\nu} + (X_n)_i \subset f - (X_n)_i + (X_n)_i \subset f + (X_n - X_n)_i$ (by Theorem 3) $\subset f + W_n$, and $\mathfrak{P}(S, g_{n\nu} + (X_n)_i) \subset \mathfrak{P}(S, f + W_n) \subset \mathfrak{P}(S, T)$. So in the sequence of open sets $g_{n\nu} + (X_n)_i, n=1, 2, \dots, \nu=1, \dots, m_n$, we can find, for every open set T and every $f \in \mathfrak{P}(S, T)$, a T' with $f \in \mathfrak{P}(S, T') \subset \mathfrak{P}(S, T)$. Therefore they form a complete system of neighborhoods in S .

Ad (8): In separable spaces regularity implies normality (cf. [5]), that

is, (8) follows from (6) and (10).†

7. We are now in a position to formulate and prove the fundamental Theorem 18.

DEFINITION 11. If D is an arbitrary set, L^D is the set of all functions $F(a)$ with the domain D and a range $\subset L$ (that is, $a \in D$, $F(a) \in L$). A function $F(a)$ is "bounded" if its range (the set of all $F(a)$, $a \in D$) is bounded (see Definition 5, this set being $\subset L$). The set of all bounded $F \in L^D$ is L_b^D . If $U \in \mathfrak{U}$, define a set $U' \subset L_b^D$ as the set of all $F \in L_b^D$ with a range $\subset U$. The set of all U' is $\mathfrak{U}'$.

Remark. If L is metric, this boundedness means that the (real-numerical) function $\|F(a)\|$ should be bounded. Our $\mathfrak{U}'$ corresponds to the metric $\|F\|'$ of L_b^D defined by $\|F\|' = \text{l.u.b.} \|F(a)\|$.

THEOREM 17. L_b^D forms with $\mathfrak{U}'$ a topological linear space, that is, it satisfies Definition 2b, (1)–(6). It is convex, that is, it also satisfies Definition 2b, (7), provided that L is convex.

All parts of this theorem are obvious.

THEOREM 18. If L is topologically complete, so is L_b^D .

Remark. It is easily seen that this statement holds for sequential completeness instead of for topological completeness. But, in view of the application to the generalized theory of almost periodic functions, and because we believe that this notion of completeness is the natural one, it is important to prove our theorem in its present form.—

Let $S^* \subset L_b^D$ be a closed and totally bounded set; we have to prove that it is compact. As every infinite $T^* \subset S^*$ contains a sequence $\mathfrak{F}_1, \mathfrak{F}_2, \dots$ of distinct elements, we may assume $T^* = (\mathfrak{F}_1, \mathfrak{F}_2, \dots)$.

For a fixed $a \in D$ denote the set of all $\mathfrak{F}(a)$, $\mathfrak{F} \in S^*$, by R_a . As $S^* \subset \mathfrak{S}(\mathfrak{F}_1 + U', \dots, \mathfrak{F}_n + U')$ implies that $R_a \subset \mathfrak{S}(\mathfrak{F}_1(a) + U, \dots, \mathfrak{F}_n(a) + U)$, R_a is totally bounded, and, with it, $R_a - R_a$ and $(R_a - R_a)_{\text{cl}}$ (by Theorems 9, 11). The latter set is also closed, and as it is contained in L , it is compact. Thus Theorem 16 applies to it; and, in particular, the open sets $(W_1)_i, (W_2)_i, \dots$ constructed in the part "Ad (9)" of its proof form a complete system of neighborhoods of 0 in $(R_a - R_a)_{\text{cl}}$. (Note that the $(W_n)_i$ do not depend on $(R_a - R_a)_{\text{cl}}$, that is, on a ; but if we choose for an open set T with $0 \in T$ the $p = p(T, a)$ with $\mathfrak{P}((R_a - R_a)_{\text{cl}}, W_p) \subset \mathfrak{P}((R_a - R_a)_{\text{cl}}, T)$ does depend on T and on a .)

Now apply the diagonal process used in the last paragraph of the proof of Theorem 15 to $\mathfrak{F}_1, \mathfrak{F}_2, \dots$ and $W'_1, W'_2, \dots$ in L_b^D (instead of $f_1, f_2, \dots$ and

† Tychonoff proves, loc. cit., only normality, that is, Hausdorff's axiom (7) ([1], p. 229). But if we replace in his result ([5], p. 140) the closed sets F, ϕ by two sets F, ϕ without common condensation points, and the space R by $F + \phi$ (in which F, ϕ are relatively closed), (8) obtains. (Hausdorff's notations for F, ϕ, R are F_1, F_2, E .)

$U_1, U_2, \dots$ in L , which we considered there). As S^* is totally bounded (as S was there) we obtain a subsequence $\mathfrak{F}^{(1)}, \mathfrak{F}^{(2)}, \dots$, such that, for each W'_p , an $n'_1 = n'_1(p)$ exists so that $m, n \geq n'_1$ imply $\mathfrak{F}^{(m)} - \mathfrak{F}^{(n)} \in W'_p$. Thus $\mathfrak{F}^{(m)}(a) - \mathfrak{F}^{(n)}(a) \in W_p$. If any open set $T \subset L$ with $0 \in T$ is given, there is a $p = p(T, a)$ with $\mathfrak{P}((R_a - R_a)_{\mathfrak{a}1}, W_p) \subset \mathfrak{P}((R_a - R_a)_{\mathfrak{a}1}, T)$. As $\mathfrak{F}^{(m)}(a) - \mathfrak{F}^{(n)}(a) \in R_a - R_a \subset (R_a - R_a)_{\mathfrak{a}1}$, we thus have $\mathfrak{F}^{(m)}(a) - \mathfrak{F}^{(n)}(a) \in T$. Putting $n_1(T, a) = n'_1(p(T, a))$, we see that $\mathfrak{F}^{(1)}(a), \mathfrak{F}^{(2)}(a), \dots$ is a fundamental sequence (in L) in the sense of Definition 9 ($m, n \geq n_1$ implying $\mathfrak{F}^{(m)}(a) - \mathfrak{F}^{(n)}(a) \in T$). As L is topologically complete, it is also sequentially complete (by Theorem 15), and so $\mathfrak{F}^{(1)}(a), \mathfrak{F}^{(2)}(a), \dots$ is convergent. Denote its limit by $\mathfrak{F}(a)$ (we know so far only that $\mathfrak{F} \in L^D$).

We constructed above an $n'_1 = n'_1(p)$, independent of a , such that, for $m, n \geq n'_1$, $\mathfrak{F}^{(m)} - \mathfrak{F}^{(n)} \in W'_p$. (Note that, for an arbitrary open set T , $0 \in T$, in the place occupied by W_p , the corresponding $n_1 = n_1(T, a)$ would have depended on a .) Thus $\mathfrak{F}^{(m)}(a) - \mathfrak{F}^{(n)}(a) \in W_p$, $\mathfrak{F}^{(m)}(a) - \mathfrak{F}(a) \in (W_p)_{\mathfrak{a}1}$.

Now assume that there exists a $U' \in \mathfrak{U}'$, such that, for infinitely many n , $\mathfrak{F}^{(n)} - \mathfrak{F}$ is not $\in U'$. Then we can select a subsequence $\mathfrak{F}_1^{(1)}, \mathfrak{F}_2^{(1)}, \dots$ from $\mathfrak{F}^{(1)}, \mathfrak{F}^{(2)}, \dots$ such that, for all n , $\mathfrak{F}_1^{(n)} - \mathfrak{F}$ is not $\in U'$. Now choose $V \in \mathfrak{U}$, $V + V \subset U$, and repeat the preceding construction with $\mathfrak{F}_1^{(1)}, \mathfrak{F}_2^{(1)}, \dots$ and $V, W_1, W_2, \dots$ instead of with $\mathfrak{F}_1, \mathfrak{F}_2, \dots$ and $W_1, W_2, \dots$. We obtain a subsequence $\mathfrak{F}_2^{(1)}, \mathfrak{F}_2^{(2)}, \dots$ of $\mathfrak{F}_1^{(1)}, \mathfrak{F}_1^{(2)}, \dots$ and a $\mathfrak{G}$ such that, for every $a \in D$, $\mathfrak{G}(a)$ is the limit of $\mathfrak{F}_2^{(1)}(a), \mathfrak{F}_2^{(2)}(a), \dots$. But this is a subsequence of $\mathfrak{F}^{(1)}(a), \mathfrak{F}^{(2)}(a), \dots$ and therefore has the limit $\mathfrak{F}(a)$; thus $\mathfrak{F}(a) = \mathfrak{G}(a)$ for all $a \in D$, and $\mathfrak{F} = \mathfrak{G}$. Furthermore, we have, from a certain n on (it is $n'_1(1)$, if $n'_1(p)$ is the analogue of $n'_1(p)$ in the present construction, and thus independent of a) $\mathfrak{F}_2^{(n)}(a) - \mathfrak{F}(a) \in V_{\mathfrak{a}1} \subset V + V \subset U$, $\mathfrak{F}_2^{(n)} - \mathfrak{F} \in U'$. But as $\mathfrak{F}_2^{(1)}, \mathfrak{F}_2^{(2)}, \dots$ is a subsequence of $\mathfrak{F}_1^{(1)}, \mathfrak{F}_1^{(2)}, \dots$ we should always have $\mathfrak{F}_2^{(n)} - \mathfrak{F}$ is not $\in U'$. This is a contradiction. Thus there is an $n'_2 = n'_2(U')$ for every $U' \in \mathfrak{U}'$ such that $n \geq n'_2$ implies $\mathfrak{F}^{(n)} - \mathfrak{F} \in U'$. (Remember that the sequence $\mathfrak{F}^{(1)}, \mathfrak{F}^{(2)}, \dots$ is independent of U' .) So if $\mathfrak{F} \in L^D$, it is the limit of the sequence $\mathfrak{F}^{(1)}, \mathfrak{F}^{(2)}, \dots$ and therefore a condensation point of the original sequence $\mathfrak{F}_1, \mathfrak{F}_2, \dots$. Thus the compactness of S^* would have been established.

Assume $U' \in \mathfrak{U}'$. We proceed as in the last paragraph of the proof of Theorem 8. Choose $V \in \mathfrak{U}$, $V + V \subset U$ and $W \in \mathfrak{U}$, $\alpha W \subset V$ if $-1 \leq \alpha \leq 1$. There is an n such that $\mathfrak{F}_n - \mathfrak{F} \in W'$, that is, all $\mathfrak{F}_n(a) - \mathfrak{F}(a) \in W$. The set of all $\mathfrak{F}_n(a)$, $a \in D$, for this n , is bounded, say $\subset \beta W$. Thus, with $\gamma = \max(1, |\beta|)$,

$$\begin{aligned} \mathfrak{F}(a) &= \mathfrak{F}_n(a) - (\mathfrak{F}_n(a) - \mathfrak{F}(a)) \in \beta W - W \\ &= \gamma \left(\frac{\beta}{\gamma} W \right) + \gamma \left(-\frac{1}{\gamma} W \right) \subset \gamma(V + V) \subset \gamma U. \end{aligned}$$

So the set of all $\mathfrak{F}(a)$, $a \in D$, is bounded, and $\mathfrak{F} \in L_b^D$. This completes the proof.

IV. NON-METRIC EXAMPLES

8. Many metric and complete linear spaces are known, so that it is not necessary to point out such examples. By Theorem 15, our notion of topological completeness gives nothing new as long as L satisfies Hausdorff's first countability axiom. For this reason those examples will be of particular interest which violate this axiom. We mentioned three such spaces in the last footnote on page 2, and we shall discuss them now in detail.

DEFINITION 12. Denote Hilbert space by $\mathfrak{H}$ and the space of all bounded linear operators in $\mathfrak{H}$ by $\mathfrak{B}$ (cf. for instance [6], Chapter I; [7], Chapter I, paragraph 2; [2], pp. 372–373). There are various ways to define sets $\mathfrak{U}$ for $L = \mathfrak{H}$ or $\mathfrak{B}$, of which we shall consider the following. (Cf. [2], pp. 378–388, where the discussion of all these topologies is to be found. The notations $\|f\|$, (f, g) , etc. are explained in each of the above references.)

(a) For any $\delta > 0$ define $U_1(\delta)$ as the set of all $f \in \mathfrak{H}$ with $\|f\| \leq \delta$; $\mathfrak{U}_1$ is the set of all $U_1(\delta)$.

(b) For any $n = 1, 2, \dots$, $\phi_1, \dots, \phi_n \in \mathfrak{H}$, $\delta > 0$, define $U_2(\phi_1, \dots, \phi_n; \delta)$ as the set of all $f \in \mathfrak{H}$ with $|(f, \phi_\nu)| \leq \delta$ for $\nu = 1, \dots, n$; $\mathfrak{U}_2$ is the set of all $U_2(\phi_1, \dots, \phi_n; \delta)$.

(c) For any $\delta > 0$ define $U_3(\delta)$ as the set of all $A \in \mathfrak{B}$ with $\|A\| \leq \delta$, where

$$\|A\| = \text{l.u.b.}_{f \in \mathfrak{H}, f \neq 0} \frac{\|Af\|}{\|f\|}$$

(that is, the set of all $A \in \mathfrak{B}$ with $\|Af\| \leq \delta \|f\|$ identically); $\mathfrak{U}_3$ is the set of all $U_3(\delta)$.

(d) For any $n = 1, 2, \dots$, $\phi_1, \dots, \phi_n \in \mathfrak{H}$, $\delta > 0$, define $U_4(\phi_1, \dots, \phi_n; \delta)$ as the set of all $A \in \mathfrak{B}$ with $\|A\phi_\nu\| \leq \delta$ for $\nu = 1, \dots, n$; $\mathfrak{U}_4$ is the set of all $U_4(\phi_1, \dots, \phi_n; \delta)$.

(e) For any $n = 1, 2, \dots$, $\phi_1, \psi_1, \dots, \phi_n, \psi_n \in \mathfrak{H}$, $\delta > 0$, define $U_5(\phi_1, \psi_1, \dots, \phi_n, \psi_n; \delta)$ as the set of all $A \in \mathfrak{B}$ with $|(A\phi_\nu, \psi_\nu)| \leq \delta$ for $\nu = 1, \dots, n$; $\mathfrak{U}_5$ is the set of all $U_5(\phi_1, \psi_1, \dots, \phi_n, \psi_n; \delta)$.

$\mathfrak{U}_1$ describes the strong, $\mathfrak{U}_2$ the weak topology of $\mathfrak{H}$; $\mathfrak{U}_3$ describes the uniform, $\mathfrak{U}_4$ the strong, $\mathfrak{U}_5$ the weak topology of $\mathfrak{B}$.

THEOREM 19. $\mathfrak{H}$ and $\mathfrak{B}$ are convex topological linear spaces (that is, they fulfill Definition 2b, (1)–(7)) in all five topologies of Definition 11.

The proof is immediate.

THEOREM 20. $\mathfrak{F}, \mathfrak{U}_1$ and $\mathfrak{B}, \mathfrak{U}_3$ are originated by metrics (in the sense of the remark after Definition 2b), with the absolute values $\|f\|$ and $\|A\|$ respectively. Both are sequentially, and thus topologically, complete.

The metric properties are verified immediately. Sequential completeness is one of the fundamental properties of $\mathfrak{F}, \mathfrak{U}_1$ (cf. [6], pp. 66 and 111), and extends from it immediately to $\mathfrak{B}, \mathfrak{U}_3$. Topological completeness follows by Theorem 15.

We now investigate the non-metric topologies $\mathfrak{F}, \mathfrak{U}_2$ and $\mathfrak{B}, \mathfrak{U}_4$ or $\mathfrak{U}_5$. Theorem 22 has some independent interest.

THEOREM 21. A set $S \subset \mathfrak{F}$ or $\mathfrak{B}$ is totally bounded in the topology $\mathfrak{U}_2$ or $\mathfrak{U}_5$ (these are the weak topologies), if and only if the (real-numerical) sets of all $|(f, \phi)|, f \in S$, or $|A\phi, \psi|, A \in S$, are bounded for every choice of $\phi \in \mathfrak{F}$ or $\phi, \psi \in \mathfrak{F}$ (no uniformity is required).

Consider $\mathfrak{F}, \mathfrak{U}_2$. If S is totally bounded, take $U = U_2(\phi; 1)$, $S \subset \mathfrak{S}(f_1 + U, \dots, f_n + U)$. Then each $f \in S$ is such that $|(f - f_\nu, \phi)| \leq 1$ for some $\nu = 1, \dots, n$, and thus

$$|(f, \phi)| = \max_{\nu=1, \dots, n} |(f_\nu, \phi)| + 1,$$

proving the necessity of our condition. To prove its sufficiency, consider a $U \in \mathfrak{U}_2$, that is, $U = U_2(\phi_1, \dots, \phi_n; \delta)$. Choose $c \geq |(f, \phi_\nu)|$ for all $f \in S$, $\nu = 1, \dots, n$, and choose $N = 1, 2, \dots$ with $2^{3/2}c/N \leq \delta$. Consider all N^{2n} systems of $2n$ integers $\rho_1, \sigma_1, \dots, \rho_n, \sigma_n = -N+1, -N+3, \dots, N-1$. If for a combination $\rho_1, \sigma_1, \dots, \rho_n, \sigma_n$ there exists an f with

$$\frac{\rho_\nu - 1}{N}c \leq \Re(f, \phi_\nu) \leq \frac{\rho_\nu + 1}{N}c,$$

$$\frac{\sigma_\nu - 1}{N}c \leq \Im(f, \phi_\nu) \leq \frac{\sigma_\nu + 1}{N}c \quad \text{for all } \nu = 1, \dots, n,$$

choose one and call it $f_{\rho_1, \sigma_1, \dots, \rho_n, \sigma_n}$. The number of these $f_{\rho_1, \sigma_1, \dots, \rho_n, \sigma_n}$ is finite, say $M \leq N^{2n}$; let us denote them by $f^{(1)}, \dots, f^{(M)}$. One easily sees that, for each $f \in S$, there exists some $f^{(\mu)}$ with

$$|\Re(f - f^{(\mu)}, \phi_\nu)| \leq \frac{2c}{N}, \quad |\Im(f - f^{(\mu)}, \phi_\nu)| \leq \frac{2c}{N} \quad \text{for all } \nu = 1, \dots, n,$$

implying that

$$|(f - f^{(\mu)}, \phi)| \leq \frac{2^{3/2}c}{N} \leq \delta.$$

Thus $S \subset \mathfrak{S}(f^{(1)} + U, \dots, f^{(M)} + U)$, proving the total boundedness of our set S , and the sufficiency of our condition.

$\mathfrak{B}$, $\mathfrak{U}_5$ are discussed in the same way.

THEOREM 22. *A set $S \subset \mathfrak{S}$ or $\mathfrak{B}$ is totally bounded in the topology $\mathfrak{U}_2$ or $\mathfrak{U}_5$, if and only if it is bounded in the topology $\mathfrak{U}_1$ or $\mathfrak{U}_3$ (these are the metric topologies), that is, if the (real-numerical) set of all $\|f\|$, $f \in S$, or $\|A\|$, $A \in S$, is bounded.*

Consider $\mathfrak{S}$, $\mathfrak{U}_2$. If $\|f\| \leq c$, then $|(f, \phi)| \leq \|f\| \cdot \|\phi\| \leq c\|\phi\|$, proving the total boundedness by Theorem 21, and thus the sufficiency of our condition. To prove the necessity, assume S totally bounded and the absolute values $\|f\|$, $f \in S$, not bounded. For each $n = 1, 2, \dots$ choose $f_n \in S$, $\|f_n\| \geq n^2$. For each ϕ , $|(f_n, \phi)|$ are bounded (by Theorem 21). Thus $\lim_{n \rightarrow \infty} \|f_n/n\| = \infty$ and, for each ϕ , $\lim_{n \rightarrow \infty} (f_n/n, \phi) = 0$, contradicting a basic property of "weak convergence" (cf. [2], p. 380, footnote 32, those considerations being a special case of a result of S. Banach, cf. [8], Theorem 5, pp. 157–160). Thus our condition is necessary.

$\mathfrak{B}$, $\mathfrak{U}_5$ are discussed in the same way (using [2], p. 382, footnote 35).

THEOREM 23. *Each of the three topologies $\mathfrak{U}_2$, for $\mathfrak{S}$, $\mathfrak{U}_4$ and $\mathfrak{U}_5$ for $\mathfrak{B}$ violate both countability axioms of Hausdorff, but all three are topologically complete.*

Hausdorff's first countability axiom is not fulfilled, by [2], p. 380 and pp. 382–383; hence the second is not fulfilled either. Every $\mathfrak{U}_2$ -totally bounded set $S \subset \mathfrak{S}$ is a subset of some set $U_1(c)$: $\|f\| \leq c$. Now the latter sets are all compact (cf. [2], p. 381, footnote 34), and S , being a closed subset of a compact set, is also compact. Therefore $\mathfrak{S}$ with $\mathfrak{U}_2$ is topologically complete. $\mathfrak{B}$ with $\mathfrak{U}_5$ is discussed in the same way.

It remains to discuss $\mathfrak{B}$ with $\mathfrak{U}_4$. We may replace every $A \in \mathfrak{B}$ by its adjoint A^* ; then $A \in \mathfrak{U}_4(\phi_1, \dots, \phi_n, \delta)$ means that $\|A^* \phi_\nu\| \leq \delta$ for $\nu = 1, \dots, n$. This means that, for every $f \in \mathfrak{S}$ with $\|f\| = 1$, $|(A^* \phi_\nu, f)| \leq \delta$ (the sufficiency follows from Schwarz's inequality, the necessity from the substitution $f = A^* \phi_\nu / \|A^* \phi_\nu\|$), that is, $|(Af, \phi_\nu)| \leq \delta$. The bounded linear operators A are characterized by the properties $A(\alpha f) = \alpha Af$ (α any complex number), $A(f+g) = Af + Ag$ within the set $\mathfrak{B}_1$ of all operators A , subject only to the restriction $A(\alpha f) = \alpha Af$ ($\alpha > 0$). The boundedness condition $|(Af, g)| \leq c \cdot \|f\| \cdot \|g\|$ (or $\|Af\| \leq c \cdot \|f\|$, cf. the remark made above), with some fixed $c > 0$, is required in any case. If we use the topology analogous to $\mathfrak{U}_4$ in $\mathfrak{B}_1$, $\mathfrak{B}$ is a closed subset of $\mathfrak{B}_1$, therefore topologically complete if $\mathfrak{B}_1$ is topologically complete. For the operators A of $\mathfrak{B}_1$ we may restrict the definition domain to the set S_1 : $\|f\| = 1$, as $A(\alpha f) = \alpha Af$ ($\alpha > 0$) then allows a unique extension to $\mathfrak{S}$. Now in this interpretation $\mathfrak{B}_1$, $\mathfrak{U}_4$ coincides exactly with the $\mathfrak{S}_1^{\mathfrak{S}_1}$ of $\mathfrak{S}$, $\mathfrak{U}_2$, from

Definition 11, and thus it is topologically complete by Theorem 18. This completes the proof.

V. THE PSEUDO-METRICS IN CONVEX SPACES

9. If L is topological and convex, that is, if L fulfills Definition 2b, (1)–(7), we can define a family of notions, each of which is an analogue to the absolute value, and which together describe the topology of L . In various applications (for instance, in the theory of almost periodic functions) this can be used to replace the metric, even when the countability axioms are violated.

THEOREM 24. *If L is topological and convex, $(\alpha U)_{01} \subset (\beta U)_i$ for $U \in \mathfrak{U}$ and $0 < \alpha < \beta$.*

By Definition 2b, (7), $U + U \subset 2U$; repeating this n times, $U + \dots + U$ (2^n times) $\subset 2^n U$, and thus, a fortiori, $(2^n - 2)U + U + U \subset 2^n U$. Therefore $(2^n - 2)U_{01} \subset (2^n - 2)U + U$ (by Theorem 5) $\subset (2^n U)_i = 2^n U_i$, $((2^n - 2)/2^n)U_{01} \subset U_i$. By Theorem 12, $0 < \alpha < 1$ implies $\alpha U_{01} \subset \alpha U_{01} + (1 - \alpha)U_{01} = U_{01}$, and, upon replacing α by α/β and multiplying by β , $0 < \alpha < \beta$, $\alpha U_{01} \subset \beta U_{01}$. Now for $0 < \alpha < \beta$ we can find an integer n for which $\alpha/\beta < (2^n - 2)/2^n < 1$, and then

$$\alpha U_{01} \subset \beta \left(\frac{2^n - 2}{2^n} U_{01} \right) \subset \beta U_i.$$

DEFINITION 13. *For $f \in L$ and $U \in \mathfrak{U}$ consider the set of all $\alpha > 0$ with $f \in \alpha U$. Its g.l.b. is denoted by $\|f\|_U^+$.*

THEOREM 25. *$\|f\|_U^+$ is finite, ≥ 0 , and continuous. If $\alpha \geq 0$, $\|\alpha f\|_U^+ = \alpha \|f\|_U^+$; furthermore, $\|f + g\|_U^+ \leq \|f\|_U^+ + \|g\|_U^+$. $\|f\|_U^+ = 0$ means that $f \in \mathfrak{P}(\alpha U \text{ over all } \alpha > 0) = (+0)U$.†*

As βf is continuous, for small $\beta > 0$, $\beta f \in U$, $f \in U/\beta$; thus the α -set is not empty and $\|f\|_U^+$ is finite; it is obviously non-negative. $\|\alpha f\|_U^+ = \alpha \|f\|_U^+$ is obvious for $\alpha > 0$; for $\alpha = 0$ it states merely that $\|0\|_U^+ = 0$. If $f \in \alpha U$, $g \in \beta U$, Theorems 12 and 24 imply that $f + g \in (\alpha + \beta)U_{01} \subset (\alpha + \beta + \delta)U$ for every $\delta > 0$; from this it follows that $\|f + g\|_U^+ \leq \|f\|_U^+ + \|g\|_U^+$. The last statement is obvious. It remains to prove the continuity of $\|f\|_U^+$.

Assume $0 < \alpha < \|f_0\|_U^+ < \beta$. (If $\|f_0\|_U^+ = 0$ then omit the α .) Choose α' , β' with $\alpha < \alpha' < \|f_0\|_U^+ < \beta' < \beta$. Then $f_0 \in \beta' U \subset \beta U_i$, f_0 is not $\in \alpha' U \supset \alpha U_{01}$, $f_0 \notin \mathfrak{P}(\beta U_i, \mathfrak{C}(\alpha U_{01}))$. If an f belongs to this set, we have $f \in \beta U_i \subset \beta U$, f is not $\in \alpha U_{01} \supset \alpha U$, and thus $\alpha \leq \|f\|_U^+ \leq \beta$. Furthermore, $\mathfrak{P}(\beta U_i, \mathfrak{C}(\alpha U_{01}))$ is obviously open, proving the continuity of $\|f\|_U^+$.

† We write $(+0)U$ in order to distinguish this set from $0U$, which of course is (0) .

THEOREM 26. *The sets $\|f - f_0\|_U^+ < \delta$, $f_0 \in L$, $U \in \mathcal{U}$, $\delta > 0$, form a complete system of neighborhoods in L . (It would be sufficient to consider $\delta = 1$.)*

As $\|f - f_0\|_U^+$ is continuous in f , these sets are all open. If S is an open set, $f_0 \in S$, there is a $U \in \mathcal{U}$ with $f_0 + U \subset S$, and $\|f - f_0\|_U^+ < 1$ implies $f - f_0 \in U$, $f \in f_0 + U \subset S$. Thus the set $\|f - f_0\|_U^+ < 1$ contains f_0 and is contained in S .

We could extend $\|\alpha f\|_U^+ = \alpha \|f\|_U^+$ from the non-negative α 's to all real α by introducing $\|f\|_U = \max(\|f\|_U^+, \|-f\|_U^+)$, but we shall not discuss this further here.

Note that in metric spaces L (cf. the remark after Definition 2b), the functions $\|f\|_U^+$ with $U = S^0(0; \delta)$ or with $U = S^1(0; \delta)$ coincide with each other and with their $\|f\|_U$, all of them being equal to $\|f\|/\delta$. If L is only topological and if we form L_b^D by Definition 11, then obviously

$$\|\mathfrak{F}\|_{U'}^+ = \text{l.u.b.}_{x \in D} \|\mathfrak{F}(x)\|_U^+.$$

APPENDIX I

10. We wish to make two remarks which are useful for some applications.

Remark 1. The linear space L , as defined in Definition 1, may allow complex numbers α (instead of the real ones alone) as factors. We then call L complex linear and we change Definition 1 so as to admit in its conditions (4), (5), (6) also complex α and β . Then Definition 2a for the metric should be formulated with complex α 's in its condition (2). But the important change is the one in Definition 2b for topological L 's: here condition (4) must include all complex α 's with $|\alpha| \leq 1$ instead of only the real α 's with $-1 \leq \alpha \leq 1$. Then Theorem 7, stating the continuity of αf as a function of α and f , can be proved without any changes. (Note that the definition of convexity, Definition 2b, (7), is unaffected, and that Definition 7, Theorems 12 and 13 remain restricted to real coefficients.)

This stronger form of Definitions 2b, (4), which we call (4'), can be replaced by the following two conditions:

(4'_1) if $U \in \mathcal{U}$, there is a $V \in \mathcal{U}$, such that, for every real α with $0 \leq \alpha \leq 1$, $\alpha V \subset U$,

(4'_2) if $U \in \mathcal{U}$, there is a $V \in \mathcal{U}$ with $iV \subset U$.

(4'_1) and (4'_2), with the help of (3) and (5), include (4'): Choose $V + V \subset U$, W with $\alpha W \subset V$ for $0 \leq \alpha \leq 1$, X with $iX \subset W$, Y with $iY \subset X$, Z with $iZ \subset Y$, $\Xi \in \mathfrak{B}(W, X, Y, Z)$, where V, W, X, Y, Z, Ξ are all $\in \mathcal{U}$. Thus $i^n \Xi \subset W$ for $n = 0, 1, 2, 3$. Now if a complex α is such that $|\alpha| \leq 1$, we can write $\alpha = i^m \beta + i^n \gamma$, $m = 0, 2$; $n = 1, 3$; β, γ real; $0 \leq \beta, \gamma \leq 1$. Thus $\alpha \Xi \subset \beta i^m \Xi + \gamma i^n \Xi \subset \beta W + \gamma W \subset V + V \subset U$, proving (4') as we stated it.

Remark 2. If L is convex and $S \subset L$, then

$$((S_{\text{conv}})_{\text{cl}})_{\text{conv}} = (S_{\text{conv}})_{\text{cl}}.$$

As the left side obviously $\supset$ the right side, we need only prove $\subset$. And as $(S_{\text{conv}})_{\text{conv}} = S_{\text{conv}}$, we may write T for S_{conv} and prove $(T_{\text{cl}})_{\text{conv}} \subset (T_{\text{conv}})_{\text{cl}}$. Assume $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}$, $V + V \subset U$. Then $V_{\text{conv}} \subset U$ (cf. the beginning of the proof of Theorem 14), and $T_{\text{cl}} \subset T + V$, $(T_{\text{cl}})_{\text{conv}} \subset T_{\text{conv}} + V_{\text{conv}} = T_{\text{conv}} + U$. Thus $(T_{\text{cl}})_{\text{conv}} \subset \mathfrak{P}(T_{\text{conv}} + U)$, where U runs over all elements of $\mathfrak{U}$. By Theorem 5, the right side is $= (T_{\text{conv}})_{\text{cl}}$, thus $(T_{\text{cl}})_{\text{conv}} \subset (T_{\text{conv}})_{\text{cl}}$, completing the proof.

APPENDIX II

11. The coefficients α , which occur in Definition 1, play an important role in the applications of this theory, but the theory itself could be developed without them. That is, we could work on the basis of Definition 1, parts (1), (2), (7) alone; in other words, the theory could be extended from linear spaces to Abelian groups (except, of course, for the statements about convexity). Definition 2a has then to be restricted to (1), (3), and Definition 2b to (1), (2), (3), (4) with $\alpha = -1$ only, and (5). Note that (1), (2), (3) are general topological axioms, while (4), (5) express that $-f$ and $f+g$ (that is, the "group operations") are continuous.

After these changes are made all our considerations remain almost unaltered, except that the notion of "boundedness" must be avoided (because Definition 2b, (6), on which it is based, has been omitted); it has to be replaced by "total boundedness."

BIBLIOGRAPHY TO "COMPLETE TOPOLOGICAL SPACES"

1. F. Hausdorff, *Mengenlehre*, Berlin, deGruyter, 1927.
2. J. von Neumann, *Zur Algebra der Funktionaloperatoren und Theorie der normalen Operatoren*, *Mathematische Annalen*, vol. 102 (1929), pp. 370-427.
3. P. Alexandroff and P. Urysohn, *Zur Theorie der topologischen Räume*, *Mathematische Annalen*, vol. 92 (1924), pp. 258-266.
4. P. Alexandroff, *Über die Struktur der bikompakten topologischen Räume*, *Mathematische Annalen*, vol. 92 (1924), pp. 267-274.
5. A. Tychonoff, *Über einen Metrisationssatz von P. Urysohn*, *Mathematische Annalen*, vol. 95 (1926), pp. 139-142.
6. J. von Neumann, *Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren*, *Mathematische Annalen*, vol. 102 (1929), pp. 49-131.
7. M. H. Stone, *Linear Transformations in Hilbert Space*, American Mathematical Society Colloquium Publications, vol. 15, New York, 1932.
8. S. Banach, *Sur les opérations dans les ensembles linéaires et leur application aux équations intégrales*, *Fundamenta Mathematicae*, vol. 3 (1922), pp. 133-181.

ALMOST PERIODIC FUNCTIONS IN GROUPS, II

The present paper is a continuation of the article by J. von Neumann on *Almost periodic functions in a group*, I [1].† Its main object is to extend the theory of almost periodicity to those functions having values which are not numbers but elements of a general linear space L . For functions of a real variable this extension was begun by Bochner [2], and then applied by him, see [3], to a problem concerning partial differential equations.

Bochner assumed L to be both complete and metric. In the present paper we shall admit more general linear spaces. We shall drop the metric but keep the completeness. Since the usual notion of completeness is based on the notion of metric, it was necessary to establish, for linear spaces, a notion of completeness independent of it. This was done in the preceding note of J. von Neumann [4]. The results of this note will be employed throughout, and we observe that, from the very beginning, we shall assume that L is linear with respect to arbitrary complex coefficients, see [4], Appendix I.

As in [1], the main difficulty to overcome was the definition and the establishment of a mean. This was done in Part I. The definition of a mean remained actually the same as in [1], but the proof of the existence of a mean necessitated a more elaborate argument, although, in broad lines, the argument does not differ essentially.

In Part II we deduce the existence and uniqueness of a Fourier expansion for any almost periodic function. It is worth pointing out that the representations occurring in the Fourier expansions of abstract almost periodic functions are the same as for numerical almost periodic functions, only the constant coefficients by which the representations are multiplied are abstract elements instead of numbers. (More than that, if in a linear manifold L different topologies are suitable for our purposes, then even the nature of the coefficients no longer determines the precise nature of abstractness of the almost periodic function.) Thus, roughly speaking, there are no more abstract almost periodic functions than numerical almost periodic functions. In particular, if a group admits of no other numerical almost periodic functions than the constant ones, there exists no non-constant abstract almost periodic function, no matter how general the range-space L may be.

† Numbers in brackets refer to the bibliography at the end of this paper.

In Part III we deduce the approximation theorem. Moreover, what is new also for numerical functions, we deduce summation theorems, that is to say, theorems concerning the construction of an almost periodic function of which only the Fourier expansion is given. For functions of Bohr these theorems were established by Bochner [5].

Finally, in Part IV we consider the special case in which L is a Hilbert space, or the space of bounded linear transformations of a Hilbert space into itself. We particularly refer to Theorems 39 and 40. Theorem 39 treats one of the rare cases of a class of almost periodic functions for which the Fourier expansions may be completely characterized by direct properties. Theorem 40 implies a necessary and sufficient criterion for a locally compact separable group to be compact. We should also mention that for the Abelian addition group of all integers, R. H. Cameron, in an unpublished paper, has found a result which has some connection with our Theorem 39.

PART I. THE MEAN-VALUE

Let L be a convex topological space which we shall assume throughout to be topologically complete (cf. [4], Definitions 1, 2b, and 10); and let $\mathfrak{G}$ denote a fixed arbitrary group. The elements of $\mathfrak{G}$ will be denoted by $a, b, c, \dots$, $x, y, z, \dots$. *Whenever a function is not specified, it will be tacitly assumed to be an element of $L_b^{\mathfrak{G}}$.* ($L_b^{\mathfrak{G}}$ is the set of all bounded functions with the domain $\mathfrak{G}$ and a range $\subset L$, cf. [4], Definition 11. $L_b^{\mathfrak{G}}$ is a topologically complete convex topological set, cf. [4], Theorem 18.)

If $F \in L_b^{\mathfrak{G}}$, the "translated" functions $r_a F = F(xa)$, $l_a F = F(ax)$ also belong to $L_b^{\mathfrak{G}}$ for any $a \in \mathfrak{G}$. We shall denote by $\mathfrak{R}_F$ the set of all $r_a F$, and by $\mathfrak{L}_F$ the set of all $l_a F$ ($a \in \mathfrak{G}$).

DEFINITION 1. *A function F is almost periodic if both sets $\mathfrak{R}_F$ and $\mathfrak{L}_F$ are totally bounded (cf. [4], Definition 6). The set of all almost periodic functions will be denoted by $\mathcal{A}p$.*

THEOREM 1. *Every constant function is almost periodic.*

The proof is obvious.

DEFINITION 2. *Given groups $\mathfrak{G}_1, \mathfrak{G}_2, \dots, \mathfrak{G}_p$; we denote by $\mathfrak{G}_1 \times \mathfrak{G}_2 \times \dots \times \mathfrak{G}_p$ the group consisting of all p -tuples*

$$x = [x_1, \dots, x_p] \quad (x_1 \in \mathfrak{G}_1, \dots, x_p \in \mathfrak{G}_p)$$

with the rules

$$\begin{aligned} [x_1, \dots, x_p][y_1, \dots, y_p] &= [x_1 y_1, \dots, x_p y_p], \\ [x_1, \dots, x_p]^{-1} &= [x_1^{-1}, \dots, x_p^{-1}]. \end{aligned}$$

And we shall say that a function is "almost periodic in $x_1, \dots, x_p$ " if it is almost periodic in $[x_1, \dots, x_p]$ in the group $\mathfrak{G}_1 \times \mathfrak{G}_2 \times \dots \times \mathfrak{G}_p$. In case $\mathfrak{G}_1 = \mathfrak{G}_2 = \dots = \mathfrak{G}_p = \mathfrak{G}$, we shall write $\mathfrak{G}^p$ for $\mathfrak{G}_1 \times \mathfrak{G}_2 \times \dots \times \mathfrak{G}_p$.

THEOREM 2. *If F is almost periodic, the set of (x, y) -functions $F_a = F(xay)$ ($\epsilon L^{\mathfrak{G}^2}$) is totally bounded ($a \in \mathfrak{G}$).*

Given $U \in \mathfrak{U}$, we choose $V \in \mathfrak{U}$, with $V + V - V \in U$. We choose elements $b_1, \dots, b_n$ such that to any y there corresponds an index $\nu = \nu(y)$ for which

$$F(zy) - F(zb_\nu) \in V \quad (z \in \mathfrak{G}).$$

(Remember for this, and for all discussions below, [4], Definition 6.) Hence

$$F(xay) - F(xab_\nu) \in V \quad (x, a \in \mathfrak{G}).$$

We consider the n functions $F_\nu(x) = F(xb_\nu)$. For each ν the set of x -functions $F_\nu(xa)$ is totally bounded in a . By a simple argument (compare [1], the corresponding part in the proof of Theorem 9) it follows that there exist elements $a_1, \dots, a_m$, and to each a there corresponds a $\mu = \mu(a)$ such that

$$F_\nu(xa) - F_\nu(xa_\mu) \in V \quad (x \in \mathfrak{G}; \nu = 1, \dots, n).$$

Hence to each a there corresponds a $\mu = \mu(a)$ such that, for all $x, y \in \mathfrak{G}$,

$$\begin{aligned} F(xay) - F(xa_\mu y) &= (F(xay) - F(xab_\nu)) + (F(xab_\nu) - F(xa_\mu b_\nu)) \\ &\quad + (F(xa_\mu b_\nu) - F(xa_\mu b)) \in V + V - V \subset U. \end{aligned}$$

THEOREM 3. *Let $F(z)$ be an almost periodic function. Let z be the product of positive or negative integer powers of elements $x', x'', \dots, x^{(p)}, a', a'', \dots, a^{(q)} \in \mathfrak{G}$ in an arbitrary fixed order, and let $F(z)$ be considered as a function*

$$(1) \quad F_{a', \dots, a^{(q)}} = F_{a', \dots, a^{(q)}}(x', \dots, x^{(p)})$$

of $L^{\mathfrak{G}^p}$, depending on the parameters $a', \dots, a^{(q)}$. The set of these functions is totally bounded in $a', \dots, a^{(q)}$.

It is easily seen that in this proof we may replace in z any set of consecutive factors which are all powers of variables or all powers of parameters by a new variable or a new parameter respectively (this may increase the indices p and q). Thus we may assume that z has the form $a'x'a''x'' \dots$ or $x'a'x''a'' \dots$. We denote the number of factors in z by k . For $k=2$, our theorem holds by Definition 1, and we are going to apply induction from k to $k+1$. Denoting the product of the first $k-1$ factors in z by ξ , we have to dispose of two cases: (i) $z = \xi xa$, (ii) $z = \xi ax$. A subscript to ξ shall indicate that special values have been assigned to all parameters occurring in ξ .

In Case i we know that to each $U \in \mathfrak{U}$ there correspond quantities $\xi_1, \dots, \xi_n$ with

$$(2) \quad F(\xi x) \in \bigotimes_{\mu=1}^n (F(\xi, x) + V').$$

(Cf. [4], Definitions 6, 11.) Given $U \in \mathfrak{A}$, let $V + V \subset U$. Replacing x by xa in (2) we get

$$(3) \quad F(\xi xa) \in \bigotimes_{\mu=1}^n (F(\xi, xa) + V').$$

We determine elements $a_1, \dots, a_m$ such that

$$(4) \quad F(za) \in \bigotimes_{\mu=1}^m (F(za_\mu) + V').$$

Putting here $z = \xi x$ and substituting the result in (3), we obtain

$$F(\xi xa) \in \bigotimes_{\mu=1}^m \bigotimes_{\nu=1}^n (F(\xi, xa_\mu) + V' + V') \subset \bigotimes_{\mu=1}^m \bigotimes_{\nu=1}^n (F(\xi, xa_\mu) + U').$$

This proves Case i.

In Case ii we determine quantities $\xi_1, \dots, \xi_n$, such that

$$F(\xi z) \in \bigotimes_{\mu=1}^n (F(\xi, z) + V');$$

hence

$$(5) \quad F(\xi ax) \in \bigotimes_{\mu=1}^n (F(\xi, ax) + V').$$

By Theorem 2 we may determine elements $a_1, \dots, a_m$ such that

$$F(yax) \in \bigotimes_{\mu=1}^m (F(ya_\mu x) + V').$$

Putting here $y = \xi$, and substituting the result in (5), we obtain

$$F(\xi ax) \in \bigotimes_{\mu=1}^m \bigotimes_{\nu=1}^n (F(\xi, a_\mu x) + U').$$

This proves Case ii.

COROLLARY. Let $F(z)$ be an almost periodic function. Let z be the product of positive or negative integer powers of variables $x', x'', \dots, x^{(p)}$, of parameters $a', a'', \dots, a^{(q)}$, and of constant elements $c', c'', \dots, c^{(r)}$, in an arbitrary fixed order. The set of functions $F(z)$ ($\in L_0^{\mathfrak{A}}$) is again totally bounded in $a', a'', \dots, a^{(q)}$.

If we consider also the elements $c', \dots, c^{(r)}$ as parameters, we get a larger set of functions $\epsilon L^{\mathfrak{G}}$ which by Theorem 3 is totally bounded. But a subset of a totally bounded set is also totally bounded.

THEOREM 4. *The function (1) of Theorem 3 is almost periodic in $x', \dots, x^{(p)}$, for any fixed values of the parameters $a', \dots, a^{(q)}$.*

If we multiply the element $[x', \dots, x^{(p)}] \epsilon \mathfrak{U}^p$ by an element $[b', \dots, b^{(p)}]$ on the right or the left, the argument z in $F(z)$ goes over into a product of positive or negative integer powers of the variables x , the constants a , and the parameters b . Replacing in the corollary of Theorem 3 the letters a and c by b and a , we find that the resulting set of functions is totally bounded in the parameters b . By Definition 1, the function (1) is almost periodic.

THEOREM 5. *If $F_1, \dots, F_k$ are almost periodic functions with values in linear spaces $L_1, \dots, L_k$ respectively, if S_κ ($\kappa=1, \dots, k$) is the range of F_κ , and if $\mathfrak{F}(f_1, \dots, f_k)$ is a function with the domain $f_\kappa \in S_\kappa$, $\kappa=1, \dots, k$, uniformly continuous in it, and with a range $\subset L$, then the function*

$$G(x) = \mathfrak{F}(F_1(x), \dots, F_k(x))$$

is almost periodic.

Definition 1 is fulfilled for the function $G(x)$ on account of [4], Theorem 9, if this theorem is applied to the function $\mathfrak{F}(F_1, \dots, F_k)$, considered as a function of the functions $F_1, \dots, F_k$, with the L_κ, L of Theorem 9 in [4] put equal to $L_{\kappa b}^{\mathfrak{G}}, L_b^{\mathfrak{G}}$, and its ranges S_κ to the $\mathfrak{R}_{F_\kappa}, \mathfrak{L}_{F_\kappa}$ respectively ($\kappa=1, \dots, k$).

THEOREM 6. *If $F, G \in Ap$, then $F \pm G \in Ap$. If $F \in Ap$, and $\alpha(x)$ is a numerical almost periodic function, then $\alpha F \in Ap$.*

If $F \in Ap$, the set of functions $r_a F = F(xa)$ is totally bounded; if we put $x=1$ we find that the range of F is totally bounded. Theorem 6 follows from Theorem 5, since the functions $f_1 + f_2, f_1 f_2$ are uniformly continuous if f_1 and f_2 run over totally bounded sets, the closure of such sets being compact and separable ([4], Theorem 11, Theorem 7, Definition 10, and Theorem 16).

THEOREM 7. *The set Ap is closed.*

Let F be a condensation point of Ap , $U \in \mathfrak{U}$, $V + V \subset U$. There is a $G \in Ap$ such that $F \in G + V'$. Obviously $r_a F \in r_a G + V'$. Choose elements $F_1, \dots, F_n \in Ap$, such that

$$\mathfrak{R}_G \subset \bigcup_{\nu=1}^n (F_\nu + V').$$

Then

$$\mathfrak{R}_F \subset \bigcup_{\nu=1}^n (F_\nu + V' + V') \subset \bigcup_{\nu=1}^n (F_\nu + U'),$$

which proves that $\mathfrak{R}_F$ is totally bounded. Similarly, $\mathfrak{X}_F$ is totally bounded.

COROLLARY. *If a sequence of almost periodic functions is uniformly convergent, the limit function is also almost periodic.*

THEOREM 8. *If $F \in Ap$, the set $S = ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$ has the following properties: (i) $S \subset Ap$, (ii) S is compact and separable, (iii) $S_{\text{conv}} = S$. The same is true also for $T = ((\mathfrak{X}_F)_{\text{conv}})_{\text{cl}}$.*

(i) follows from Theorems 4, 6 and 7. (ii) follows from the fact that S is totally bounded ([4], Theorems 14 and 16). (iii) was proved in [4] Appendix I.

THEOREM 9. *If $G \in ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$, then*

$$(6) \quad ((\mathfrak{R}_G)_{\text{conv}})_{\text{cl}} \subset ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}.$$

The same is true if we replace $\mathfrak{R}$ by $\mathfrak{X}$.

For any $U \in \mathfrak{U}$, we have by assumption ([4], Theorem 5): $G \subset (\mathfrak{R}_F)_{\text{conv}} + U'$. This immediately leads to the result that $\mathfrak{R}_G \subset (\mathfrak{R}_F)_{\text{conv}} + U'$. Since U is arbitrary, it follows ([4], Theorem 5) that $\mathfrak{R}_G \subset ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$. The final relation (6) is now obvious by Theorem 8, (iii).

DEFINITION 3. *If two almost periodic functions G, F are in the relation (6), we shall write*

$$G \rightarrow F.$$

Remark. It follows from Theorem 9 that the relation $G \rightarrow F$ is equivalent to $G \in ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$.

THEOREM 10. *$F \rightarrow G$ and $G \rightarrow H$ imply $F \rightarrow H$.*

This follows immediately from Theorem 9.

THEOREM 11. *If $F \in Ap$, and $U \in \mathfrak{U}$, there is a $G \rightarrow F$, and a number $\mu \geq 0$, such that for $H \rightarrow G$ and $a \in \mathfrak{G}$ (cf. [4], Definition 13),*

$$\|H(a)\|_U^+ = \mu.$$

We consider, for $K \in ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$, the numerical function

$$(7) \quad \|K\|_{U'}^+ = \text{l.u.b.}_{x \in \mathfrak{G}} \|K(x)\|_{U'}^+.$$

It is a continuous function; hence it assumes, on every compact separable set, a maximum and a minimum; and it is easily seen that, for $K_1 \rightarrow K_2$, $\|K_1\|_{U'}^+ \leq \|K_2\|_{U'}^+$. Therefore, if G is an element for which (7) attains its minimum, and if μ denotes the minimum value, we have

$$\|K\|_{U'}^+ = \|H\|_{U'}^+ = \mu, \quad \text{for } K \rightarrow H \rightarrow G.$$

Thus, for $a_1, \dots, a_n \in \mathfrak{G}; \alpha_1, \dots, \alpha_n \geq 0; \alpha_1 + \dots + \alpha_n = 1$,

$$\text{l.u.b.}_{x \in \mathfrak{G}} \|\alpha_1 H(a_1 x) + \dots + \alpha_n H(a_n x)\|_U^+ = \text{l.u.b.}_{x \in \mathfrak{G}} \|H(x)\|_U^+ = \mu.$$

Connecting this with

$$\mu \geq \text{l.u.b.}_{x \in \mathfrak{G}} \|H(x)\|_U^+ \geq \text{l.u.b.}_{x \in \mathfrak{G}} \sum_{r=1}^n \alpha_r \|H(a_r x)\|_U^+,$$

we find for the numerical function $\gamma(x) = \|H(x)\|_U^+$ the property

$$(8) \quad \text{l.u.b.}_{x \in \mathfrak{G}} \gamma(x) = \text{l.u.b.}_{x \in \mathfrak{G}} (\alpha_1 \gamma(a_1 x) + \dots + \alpha_n \gamma(a_n x)).$$

By Theorem 5, $\gamma(x)$ is almost periodic; and (8) proves that

$$M_x \gamma(x) = \text{l.u.b.}_{x \in \mathfrak{G}} \gamma(x).$$

(Cf. the definition of M_x in [1], Definition 4.) Hence $\gamma(x)$ is constant (cf. [1], Theorem 7, (4), putting $f(x) = [\text{l.u.b.}_{x \in \mathfrak{G}} \gamma(x)] - \gamma(x)$), and the proof of our theorem is completed.

COROLLARY. *If $F \in \mathcal{A}p$, $U \in \mathcal{U}$, $f \in L$, there is a $G \rightarrow F$, and a number $\mu \geq 0$, such that for $H \rightarrow G$, $a \in \mathfrak{G}$,*

$$(9) \quad \|H(a) - f\|_U^+ = \mu.$$

Apply Theorem 11 to $F_1(x) = F(x) - f$, and denote the resulting $G_1(x)$ by $G(x) - f$.

THEOREM 12. *If $F \in \mathcal{A}p$, there is an $H \rightarrow F$ which is constant.*

We denote by S the range of $F(x)$, and by T the set $(S_{\text{con} \nabla})_{\text{cl}}$. S is totally bounded (compare the proof of Theorem 6). Hence by [4], Theorems 11, 14, T is totally bounded too. If $f \in L$ is a value of a function $G \rightarrow F$, then f is a condensation point of elements of the form

$$\alpha_1 F(a_1 x) + \dots + \alpha_n F(a_n x), \quad \alpha_1, \dots, \alpha_n \geq 0; \alpha_1 + \dots + \alpha_n = 1;$$

but an element of this form is contained in $S_{\text{con} \nabla}$; thus $f \in T$. Therefore there exists a compact separable set $T \subset L$ which contains the ranges of all $G \rightarrow F$.

Let $f_1, f_2, \dots$ be a dense sequence in T , $W_1, W_2, W_3, \dots$ a complete set of open neighborhoods of zero. Write all pairs $n, p = 1, 2, \dots$ in a sequence $n_k, p_k, k = 1, 2, \dots$, and define a sequence of functions $G_0, G_1, G_2, \dots$ in this manner: $G_0 = F$; G_{k+1} is the G of the corollary to Theorem 11 if applied to $F = G_k$, $U = U_k = W_{p_k}$, $f = f_{n_k}$. Thus

$$F \rightarrow G_0 \rightarrow G_1 \rightarrow G_2 \rightarrow \dots,$$

and for all $H \rightarrow G_k$,

$$(10) \quad \|H(x) - f_{n_k}\|_{U_k}^+ = \mu_k,$$

where μ_k is a number depending on k .

The sets $((\mathfrak{R}_{G_k})_{\text{conv}})_{01}$ are monotonely decreasing as $k \rightarrow \infty$, all closed and non-empty, and all subsets of the compact separable set $((\mathfrak{R}_F)_{\text{conv}})_{01}$; thus they have a common element H . Relation (10) means that for any p and any $x, y \in \mathfrak{G}$, the relation

$$(11) \quad \|H(x) - f\|_{W_p}^+ = \|H(y) - f\|_{W_p}^+$$

holds for all $f = f_n$. The f_n being dense, and the pseudo-metric being continuous, the relation (11) holds for all $f \in T$. Putting $f = H(y)$ we obtain $\|H(x) - H(y)\|_{W_p}^+ = 0$; hence $H(x) - H(y) \in W_p$ ($p = 1, 2, \dots$) (cf. [4], Definition 13); hence $H(x) = H(y)$.

COROLLARY. *If $F \in Ap$, there is a $\phi \in L$, and a sequence of systems ($n = 1, 2, 3, \dots$)*

$$(12) \quad \alpha_{n,1}, \dots, \alpha_{n,m_n} \geq 0, \alpha_{n,1} + \dots + \alpha_{n,m_n} = 1, a_{n,1}, \dots, a_{n,m_n} \in \mathfrak{G}$$

such that for every $U \in \mathfrak{U}$ there exists an $n_1 = n_1(U)$, for which $n \geq n_1$ implies

$$\alpha_{n,1}F(a_{n,1}x) + \dots + \alpha_{n,m_n}F(a_{n,m_n}x) - \phi \in U.$$

We denote the constant value of the function H of Theorem 12 by ϕ . Since $(\mathfrak{R}_F)_{\text{conv}}$ is separable, and H is an element of its closure, H is the limit of a sequence of (equal or different) elements of $(\mathfrak{R}_F)_{\text{conv}}$. Hence the corollary.

THEOREM 13. *If $F \in Ap$, there is a $\psi \in L$, and a sequence of systems (12), such that for every $U \in \mathfrak{U}$ there exists an $n_1 = n_1(U)$, for which $n \geq n_1$ implies*

$$\alpha_{n,1}F(xa_{n,1}y) + \dots + \alpha_{n,m_n}F(xa_{n,m_n}y) - \psi \in U.$$

Apply the corollary to Theorem 12 to the function $F(z^{-1}y)$ of $L^{\mathfrak{G}^2}$ (cf. Definition 2), and then replace z^{-1} by x .

DEFINITION 4. *If $F \in Ap$, every $\psi \in L$ which has the property described in Theorem 13 is a mean of F .*

THEOREM 14. *If $F, G \in Ap$, and ψ, χ are respective means of F, G , then $\psi \pm \chi$ is a mean of $F \pm G$.*

Given $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}$, with $2V - 2V \in U$. Suppose that $\alpha_\mu \geq 0$, $\alpha_1 + \dots + \alpha_m = 1$, $a_\mu \in \mathfrak{G}$, $\beta_\nu \geq 0$, $\beta_1 + \dots + \beta_n = 1$, $b_\nu \in \mathfrak{G}$, and that

$$(13) \quad \alpha_1 F(xa_1y) + \dots + \alpha_m F(xa_my) - \psi \in V,$$

$$(14) \quad \beta_1 F(xb_1y) + \dots + \beta_n F(xb_ny) - \chi \in V.$$

In (13) we replace y by b, y , multiply the equation by β_v , and sum over v . We obtain, using [4], Theorem 12,

$$\sum_{\mu=1}^m \sum_{v=1}^n \alpha_{\mu} \beta_v F(x a_{\mu} b_v y) - \psi \epsilon \sum_{v=1}^n \beta_v V \subset 2V.$$

Similarly, if in (14) we replace x by $x a_{\mu}$, multiply by α_{μ} and sum over μ , we get

$$\sum_{\mu=1}^m \sum_{v=1}^n \alpha_{\mu} \beta_v G(x a_{\mu} b_v y) - \chi \epsilon \sum_{\mu=1}^m \alpha_{\mu} V \subset 2V.$$

Hence

$$\sum_{\mu=1}^m \sum_{v=1}^n \alpha_{\mu} \beta_v (F(x a_{\mu} b_v y) - G(x a_{\mu} b_v y)) - (\psi - \chi) \epsilon 2V - 2V \subset U.$$

As

$$\alpha_{\mu} \beta_v \geq 0, \quad \sum_{\mu=1}^m \sum_{v=1}^n \alpha_{\mu} \beta_v \geq 1, \quad a_{\mu} b_v \in \mathfrak{G},$$

and U was arbitrary, we conclude that $\psi - \chi$ is a mean of $F - G$. A similar argument proves that $\psi + \chi$ is a mean of $F + G$.

COROLLARY. *If $F_1, F_2, \dots, F_N \in \mathcal{A}p$ ($N = 1, 2, 3, \dots$), with the means $\psi_1, \psi_2, \dots, \psi_N$ respectively, and if $U \in \mathfrak{U}$, then there exist numbers $\alpha_1, \dots, \alpha_n \geq 0$, $\alpha_1 + \dots + \alpha_n = 1$, and elements $a_1, \dots, a_n \in \mathfrak{G}$, such that simultaneously for $v = 1, 2, \dots, N$,*

$$\alpha_1 F_v(x a_1 y) + \alpha_2 F_v(x a_2 y) + \dots + \alpha_n F_v(x a_n y) - \psi_v \subset U.$$

The case $N = 2$ was treated in the proof of Theorem 14, and the same argument can be used to extend our statement from N to $N + 1$.

THEOREM 15. *Every almost periodic function has one and only one mean.*

If ψ and χ are both means of F , $\psi - \chi$ is a mean of $F - F = 0$. But every mean of 0 is 0. Hence the uniqueness of the mean.

DEFINITION 5. *If $F \in \mathcal{A}p$, its (unique) mean will be denoted by MF or $M_x F(x)$.*

THEOREM 16. *The mean has the following properties:*

1. *If $F(x) = \phi$ (constant $\in L$), then $MF = \phi$.*
2. *If α is a number, $M(\alpha F) = \alpha MF$.*
3. *$M(F \pm G) = MF \pm MG$.*
4. *$M_x F(ax) = M_x F(x)$.*
5. *$M_x F(xa) = M_x F(x)$.*
6. *$M_x F(x^{-1}) = M_x F(x)$.*
7. *If $U \in \mathfrak{U}$, $\|MF\|_U^+ \leq M_x(\|F(x)\|_U^+) \leq \|F\|_U^+$.*
8. *If $U \in \mathfrak{U}$, and $F - G \in U$, then $MF - MG \in 2U$.*

1 and 2 are obvious. 3 follows from Theorem 14. 4, 5, 6 are easily deduced from Definition 5. The first half of 7 follows from the relation

$$\|\alpha_1 F(xa_1y) + \cdots + \alpha_n F(xa_ny)\|_U^+ \leq \alpha_1 \|F(xa_1y)\|_U^+ + \cdots + \alpha_n \|F(xa_ny)\|_U^+$$

and the continuity of the pseudo-metric; the second half is a proved theorem on real almost periodic functions (cf. [1], Theorem 7). With regard to 8, from $\|F - G\|_U^+ \leq 1$, it follows that $\|MF - MG\|_U^+ \leq 1$, and thus $MF - MG \in 2U$.

THEOREM 17. *The properties 1, 2, 3, 4 (or 5), 8, of Theorem 16, determine our mean uniquely.*

Replacing ab by ba in $\mathfrak{G}$ shows that it is sufficient to consider the properties 1, 2, 3, 4, 8. Let NF be a notion defined in $A\mathfrak{p}$ satisfying these properties. If $V + V \subset U \in \mathfrak{U}$, $\alpha_1 + \cdots + \alpha_n = 1$, and

$$\alpha_1 F(a_1x) + \cdots + \alpha_n F(a_nx) - MF \subset V,$$

then

$$NF - MF = N(\alpha_1 F(a_1x) + \cdots + \alpha_n F(a_nx) - MF) \subset 2V \subset U.$$

U being arbitrary, $NF = MF$.

Remark. The properties 3, 7 of Theorem 16 imply the property 8. Hence we have that the properties 1, 2, 3, 4 (or 5), 7, of Theorem 16 determine our mean uniquely.

THEOREM 18. *If $x \in \mathfrak{G}$, $y \in \mathfrak{G}'$, and $F(x, y)$ is almost periodic in x, y , then $F(x, y)$ is almost periodic in x for fixed y , and almost periodic in y for fixed x , $G(x) = M_y F(x, y)$ is almost periodic in x , $H(y) = M_x F(x, y)$ is almost periodic in y , and*

$$(15) \quad \begin{aligned} MF &= M_x G(x) = M_x (M_y F(x, y)), \\ MF &= M_y H(y) = M_y (M_x F(x, y)). \end{aligned}$$

If the sets $F(ax, by)$, $F(xa, yb)$ are totally bounded for $a \in \mathfrak{G}$, $b \in \mathfrak{G}'$, then the sets $F(ax, y)$, $F(xa, y)$, which are parts of them, are totally bounded for $a \in \mathfrak{G}$. Hence $F(x, y)$ is almost periodic in x for fixed y . Similarly, if we interchange x and y , $F(x, y)$ is almost periodic in y for fixed x .

Given $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}$, with $V + V \subset U$, and elements $a_1, \dots, a_m, b_1, \dots, b_n$, such that

$$\begin{aligned} F(ax, y) &\in \mathfrak{S}(F(a_1x, y) + V', \dots, F(a_mx, y) + V'), \\ F(xb, y) &\in \mathfrak{S}(F(xb_1, y) + V', \dots, F(xb_n, y) + V'). \end{aligned}$$

By Theorem 16, 8,

$$G(ax) = M_y F(ax, y) \in \mathfrak{S}(M_y F(a_1x, y) + U', \dots, M_y F(a_mx, y) + U'),$$

and similarly

$$G(xb)\epsilon\mathfrak{S}(M_{\mathbf{y}}F(xb_1, y) + U', \dots, M_{\mathbf{y}}F(xb_n, y) + U').$$

Hence $G(x)$ is almost periodic, and similarly $H(y)$. In order to prove the first half of (15) we have only to show that the "mean" $NF = M_z(M_{\mathbf{y}}F)$ satisfies the conditions 1, 2, 3, 4, 7, of Theorem 16. But this is easily verified; for instance, in the case of the first half of condition 7, we have

$$\|NF\|_U^+ = \|M_z(M_{\mathbf{y}}F(x, y))\|_U^+ \leq M_z(\|M_{\mathbf{y}}F(x, y)\|_U^+) \leq M_z M_{\mathbf{y}}(\|F\|_U^+) = N(\|F\|_U^+)$$

(cf. [1], Theorem 10). Similarly for the second half of (15).

PART II. FOURIER EXPANSIONS

DEFINITION 6. If $\phi(x)$ is a numerical almost periodic function, and $F \in Ap$, $\phi \times F$ is the function

$$M_{\mathbf{y}}(\phi(xy^{-1})F(y)) = M_{\mathbf{y}}(\phi(xy)F(y^{-1}))$$

which again belongs to Ap .

THEOREM 19. $\phi \times F$ is linear in ϕ and in F , and associative: $(\phi \times \psi) \times F = \phi \times (\psi \times F)$.

The first follows from Theorem 16, 2, 3. The second follows from the following formal calculations, each step of which is justified by one of the foregoing theorems:

$$\begin{aligned} (\phi \times \psi) \times F &= M_{\mathbf{y}}(M_z(\phi(xy^{-1}z^{-1})\psi(z))F(y)) = M_{\mathbf{y}}M_z(\phi(xy^{-1}z^{-1})\psi(z)F(y)) \\ &= M_{(\mathbf{y}, z)}(\phi(xy^{-1}z^{-1})\psi(z)F(y)). \end{aligned}$$

$$\begin{aligned} \phi \times (\psi \times F) &= M_{\mathbf{y}}(\phi(xy^{-1})M_z(\psi(yz^{-1})F(z))) = M_{\mathbf{y}}M_z(\phi(xy^{-1})\psi(yz^{-1})F(z)) \\ &= M_z M_{\mathbf{y}}(\phi(xy^{-1})\psi(yz^{-1})F(z)) = M_z(M_{\mathbf{y}}(\phi(xy^{-1})\psi(yz^{-1}))F(z)), \end{aligned}$$

and substituting yz for y ,

$$\begin{aligned} &= M_z(M_{\mathbf{y}}(\phi(xz^{-1}y^{-1})\psi(y))F(z)) = M_z M_{\mathbf{y}}(\phi(xz^{-1}y^{-1})\psi(y)F(z)) \\ &= M_{(\mathbf{y}, z)}(\phi(xy^{-1}z^{-1})\psi(z)F(y)). \end{aligned}$$

THEOREM 20. If $F(x, y)$ is an almost periodic function of $L^{\mathfrak{G} \times \mathfrak{G}'}$, and $U \in \mathfrak{U}$, there exist numbers $\alpha_1, \dots, \alpha_n \geq 0$, $\alpha_1 + \dots + \alpha_n = 1$, and elements $a_r \in \mathfrak{G}$, such that, for all $x \in \mathfrak{G}$, $y \in \mathfrak{G}'$,

$$(16) \quad \sum_{r=1}^n \alpha_r F(a_r x, y) - M_z F(x, y) \epsilon U.$$

For fixed y , $F(x, y)$ is almost periodic in x . Hence, by the corollary to Theorem 14, if any finite number of fixed values is assigned to the variable y , it is possible to choose quantities α_r , a_r , which satisfy (16) for all x .

Now, choose $V \in \mathfrak{U}$, with $2V + V - 2V \subset U$, and then elements $b_1, \dots, b_m$, such that $F(x, yb) \in \mathfrak{S}(F(x, yb_1) + V', \dots, F(x, yb_m) + V')$. Putting $y = 1$, we obtain

$$F(x, b) \in \mathfrak{S}(F(x, b_1) + V', \dots, F(x, b_m) + V').$$

We determine quantities α_r, a_r such that

$$(17) \quad \sum_{r=1}^n \alpha_r F(a_r x, b_\mu) - M_x F(x, b_\mu) \in V$$

for $\mu = 1, \dots, m$. To each b there corresponds a b_μ such that

$$(18) \quad F(x, b) - F(x, b_\mu) \in V.$$

Hence by Theorem 16, 8,

$$(19) \quad M_x F(x, b) - M_x F(x, b_\mu) \in 2V.$$

On the other hand, it follows from (18) that

$$(20) \quad \sum_{r=1}^n \alpha_r F(a_r x, b) - \sum_{r=1}^n \alpha_r F(a_r x, b_\mu) \in 2V.$$

Combining (17), (19), (20), we obtain, for any $b \in \mathfrak{G}'$,

$$\sum_{r=1}^n \alpha_r F(a_r x, b) - M_x F(x, b) \in 2V + V - 2V \subset U,$$

and this proves the theorem.

DEFINITION 7. A weight function is a real almost periodic function $\phi(x)$ with the properties: $\phi(x) \geq 0$, $M\phi = 1$. A special weight function is a weight function which is a finite linear aggregate of representation coefficients. (Cf. [1], Chapter III.)

THEOREM 21. If $F \in \mathcal{A}p$, and $\phi(x)$ is a weight function, then $\phi \times F \in \mathcal{A}p$.

If $U \in \mathfrak{U}$, choose $V \in \mathfrak{U}$, with $V + V \subset U$. Then $\phi \times F = M_y(\phi(y)F(y^{-1}x))$. By Theorem 20 there are numbers $\alpha_1, \dots, \alpha_n \geq 0$, $\alpha_1 + \dots + \alpha_n = 1$, and elements $a_1, \dots, a_n \in \mathfrak{G}$, such that

$$(21) \quad \phi \times F - \sum_{r=1}^n \alpha_r \phi(a_r y) F(y^{-1} a_r^{-1} x) \in V.$$

The function

$$\psi(y) = \sum_{r=1}^n \alpha_r \phi(a_r y)$$

is a weight function. Hence there exist elements y_1, y_2 , such that $\psi(y_1) \geq 1$, $\psi(y_2) \leq 1$. Therefore we may find a γ , $0 \leq \gamma \leq 1$, with $\gamma\psi(y_1) + (1-\gamma)\psi(y_2) = 1$. From (20) we easily deduce that

$$\phi \times F \subset \gamma \sum_{r=1}^n \alpha_r \phi(a_r y_1) F(y_1^{-1} a_r^{-1} x) + (1-\gamma) \sum_{r=1}^n \alpha_r \phi(a_r y_2) F(y_2^{-1} a_r^{-1} x) + U'.$$

It is easily seen that the function in x on the right side $\epsilon(\mathfrak{R}_F)_{\text{conv}}$. Hence $\phi \times F \epsilon(\mathfrak{R}_F)_{\text{conv}} + U'$, for any $U' \in \mathfrak{U}'$. Thus ([4], Theorem 5) $\phi \times F \subset ((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$.

LEMMA 1. *If $\phi(x, y)$ is a real almost periodic function in x, y , then*

$$\psi(x) = \text{l.u.b.}_{y \in \mathfrak{G}'} \phi(x, y)$$

is almost periodic in x .

If $|\phi(ax, by) - \phi(a, x, b, y)| \leq \epsilon$, then also

$$\left| \text{l.u.b.}_{y \in \mathfrak{G}'} \phi(ax, by) - \text{l.u.b.}_{y \in \mathfrak{G}'} \phi(a, x, b, y) \right| \leq \epsilon,$$

that is,

$$\left| \text{l.u.b.}_{y \in \mathfrak{G}'} \phi(ax, y) - \text{l.u.b.}_{y \in \mathfrak{G}'} \phi(a, x, y) \right| \leq \epsilon.$$

Hence, if $\mathfrak{R}_\phi$ is totally bounded, then so is $\mathfrak{R}_\psi$. Similarly for $\mathfrak{R}_\psi$.

LEMMA 2. *If $\psi(x)$ is a weight function, and $\epsilon > 0$, there is a special weight function $\chi(x)$ such that*

$$|\psi(x) - \chi(x)| \leq \epsilon, \quad x \in \mathfrak{G}.$$

By the approximation theorem ([1], Theorem 30), for each $\delta > 0$, there is a finite linear aggregate of representation coefficients $\chi_1(x)$ such that

$$(22) \quad |\psi(x) - \chi_1(x)| \leq \delta.$$

We may assume $\chi_1(x)$ real, otherwise we replace it by $(\chi_1(x) + \overline{\chi_1(x)})/2$. After that, we may assume it ≥ 0 , otherwise $\chi_1(x) + \delta$ satisfies (22) with 2δ instead of δ . Since $M\psi = 1$, (22) implies $1 - \delta \leq M\chi_1 \leq 1 + \delta$; hence, for $\chi = \chi_1/(M\chi_1)$, we obtain

$$\begin{aligned} |\psi(x) - \chi(x)| &\leq |\psi(x) - \chi_1(x)| + \left| \left(1 - \frac{1}{M\chi_1}\right) \chi_1 \right| \\ &\leq \delta + \frac{\delta}{1 - \delta} \left[\text{l.u.b.}_{x \in \mathfrak{G}} \psi(x) + \delta \right], \end{aligned}$$

and this is $\leq \epsilon$, if δ is small enough.

THEOREM 22. *If $F \in Ap$, and $U \in \mathfrak{U}$, there exists a special weight function χ such that $\chi \times F - F \in U'$.*

$F(x^{-1}y) - F(y)$ is almost periodic in x, y , hence $\|F(x^{-1}y) - F(y)\|_U^+$ is also almost periodic in x, y (Theorem 5). By Lemma 1,

$$t(x) = \text{l.u.b.}_{y \in \mathfrak{G}} \|F(x^{-1}y) - F(y)\|_U^+$$

is almost periodic in x .

Given $\epsilon > 0$, we form with the function

$$\omega_\epsilon(u) = \begin{cases} 1 - \frac{|u|}{\epsilon} & \text{for } |u| \leq \epsilon, \\ 0 & \text{for } |u| \geq \epsilon \end{cases}$$

the almost periodic (Theorem 5) function $\psi_\epsilon(x) = \omega_\epsilon(t(x))$. We have (1) $\psi_\epsilon(x) \geq 0$, (2) $\psi_\epsilon(1) = 1$, since $t(0) = 0$, (3) if $\psi_\epsilon(x) > 0$, then $\|F(x^{-1}y) - F(y)\|_U^+ \leq \epsilon$ for all $y \in \mathfrak{G}$, (4) $m_\epsilon = M_x \psi_\epsilon(x) > 0$ by [1], Theorem 7, 4. Therefore

$$\|\psi_\epsilon(z)(F(z^{-1}x) - F(x))\|_U^+ \leq \psi_\epsilon(z) \|F(z^{-1}x) - F(x)\|_U^+ \leq \psi_\epsilon(z) t(z),$$

and this is equal to 0 if $\psi_\epsilon(z) = 0$, and $\leq \epsilon \psi_\epsilon(z)$ if $\psi_\epsilon(z) \neq 0$. Thus

$$\|\psi_\epsilon \times F - m_\epsilon F\|_U^+ \leq \epsilon m_\epsilon,$$

and for the weight function $\psi(x) = \psi_\epsilon(x)/m_\epsilon$ we have

$$\|\psi \times F - F\|_U^+ \leq \epsilon.$$

For the special weight function $\chi(x)$ from Lemma 2 we obtain

$$\begin{aligned} \|\chi \times F - F\|_U^+ &\leq \epsilon + \|(\chi - \psi) \times F\|_U^+ \leq \epsilon + M_x \|(\psi(xy^{-1}) - \chi(xy^{-1}))F(y)\|_U^+ \\ &\leq \epsilon + \epsilon M_x (\max \| \pm F(y) \|_U^+) \leq \epsilon + \epsilon C, \end{aligned}$$

where C is a constant independent of ϵ . Choosing $\epsilon < (1+C)^{-1}$, we obtain the theorem.

THEOREM 23. *If $F \in Ap$, there exists a sequence of special weight functions $\phi_n(x)$, $n=1, 2, 3, \dots$, such that F is the limit of the sequence $\phi_n \times F$, $n=1, 2, 3, \dots$.*

Consider the functions $\phi \times F$, for all special weight functions ϕ , and denote their set by S . By Theorem 22, $F \in S_{01}$; but S_{01} , being $\subset ((\mathfrak{R}_F)_{\text{nonv}})_{01}$ (Theorem

21), is separable; hence, F is the limit of a sequence of (equal or different) elements ϵS .

DEFINITION 8. If $D(x) = \{D_{\rho\sigma}(x)\}_{\rho,\sigma=1,\dots,s}$ is an irreducible normal representation of $\mathfrak{G}$ (cf. [1], Chapter III, Definitions 9, 10) we form, for $F \in Ap$,

$$f_{\rho\sigma}(D) = M_x(D_{\rho\sigma}(x^{-1})F(x)) = M(\overline{D_{\rho\sigma}F}),$$

and the matrix

$$f^D = \{f_{\rho\sigma}(D)\}_{\rho,\sigma=1,\dots,s}.$$

$f_{\rho\sigma}(D)$ is the (D, ρ, σ) -expansion-coefficient of F , f^D its D -expansion-matrix.

THEOREM 24. The expansion coefficients and matrices have the following properties:

1. $M(\overline{D_{\rho\sigma}(\alpha F)}) = \alpha M(\overline{D_{\rho\sigma}F})$, $\alpha \in L$.
2. $M(\overline{D_{\rho\sigma}(F+G)}) = M(\overline{D_{\rho\sigma}F}) + M(\overline{D_{\rho\sigma}G})$.
3. If F is the limit of a sequence F_N , $N=1, 2, \dots$, then, for each (D, ρ, σ) , $M(\overline{D_{\rho\sigma}F})$ is the limit of the sequence $M(\overline{D_{\rho\sigma}F_N})$, $N=1, 2, \dots$.
4. If $G = \phi \times F$, and if $g_{\rho\sigma}(D)$, $\phi_{\rho\sigma}(D)$, $f_{\rho\sigma}(D)$ are expansion coefficients of G , ϕ , F respectively, then

$$g^D = \phi^D f^D,$$

that is to say,

$$g_{\rho\sigma}(D) = \sum_{\tau=1}^s \phi_{\rho\tau}(D) f_{\tau\sigma}(D).$$

In particular, the matrix g^D vanishes if f^D or ϕ^D vanishes (or both).

5. If D^N , $N=1, 2, 3, \dots$, is a (finite or countable) sequence of irreducible inequivalent normal representations, if the s^N are their respective degrees, if $h_{\rho\sigma}$ are elements of L , and if

$$(23) \quad F = \sum_{N=1}^{\infty} \left(s^N \sum_{\rho,\sigma=1}^{s^N} h_{\rho\sigma}^N D_{\rho\sigma}^N \right)$$

[meaning that $F \in Ap$ is the limit, $m \rightarrow \infty$, of the sequence of elements

$$F_m = \sum_{N=1}^m \left(s^N \sum_{\rho,\sigma=1}^{s^N} h_{\rho\sigma}^N D_{\rho\sigma}^N \right)$$

from Ap], then

$$M(\overline{D_{\rho\sigma}F}) = \begin{cases} 0 & \text{if } D \neq D^N \\ h_{\rho\sigma}^N & \text{if } D = D^N \end{cases} \quad (N = 1, 2, 3, \dots).$$

1 and 2 are obvious; 3 is easily deducible from the fact that $G(x, y) \in U$ implies $M_y G(x, y) \in 2U$ (Theorem 16, 8); and 4 follows from the relation

$$\begin{aligned} g_{\rho\sigma}^D &= M_z(D_{\sigma\rho}(x^{-1})M_y(\phi(xy^{-1})F(y))) = M_y(F(y)M_z(D_{\sigma\rho}(x^{-1})\phi(xy^{-1}))) \\ &= M_y(F(y)M_z(D_{\sigma\rho}(y^{-1}z^{-1})\phi(z))) \\ &= \sum_{\tau=1}^n M_y(F(y)D_{\sigma\tau}(y^{-1})) \cdot M_z(D_{\tau\rho}(z^{-1})\phi(z)). \end{aligned}$$

As regards 5, if the number of the representations D^N is finite, it follows by direct computation of $M_z(\overline{D_{\rho\sigma}}(x)F(x))$, applying 1 and 2 and [1], Theorem 21, and in the general case by using this and 3.

THEOREM 25. *If $F \in Ap$, it has only countably many expansion matrices $\neq 0$.*

By Theorem 24, 4 and 3, this is true if F has the form $\phi \times G$, with ϕ a special weight function, or if F is the limit of such functions. Now apply Theorem 23.

DEFINITION 9. *If $F \in Ap$; if $D^N, N=1, 2, \dots$, is the sequence of irreducible normal representations (in any fixed order) for which its expansion matrices $\neq 0$; if the s^N are their respective degrees; and if $f_{N,\rho\sigma} = M_z(D_{\rho\sigma}^N(x^{-1})F(x))$; then we call the formal series*

$$(24) \quad \sum_{N=1}^{\infty} s^N \sum_{\rho, \sigma=1}^{s^N} f_{N,\rho\sigma} D_{\rho\sigma}^N$$

the Fourier expansion of F .

We call a sequence $F_m, m=1, 2, \dots$, formally convergent (to F), if the sequence of the Fourier expansions of the functions F_m is formally convergent (to the Fourier expansion of F), i.e., if for any (D, ρ, σ) , the sequence $M(\overline{D_{\rho\sigma}}F_m)$ has a limit (namely $M(\overline{D_{\rho\sigma}}F)$).

Remark. Theorem 24 states properties of the Fourier expansion. 1 and 2 state additivity. 3 states that the Fourier series of a limit is the formal limit of the Fourier series. 4 gives an important rule for the computation of the Fourier series of a convolution of an almost periodic function with a numerical almost periodic function. Finally, 5 states that the sum of a uniformly convergent series of the form (23) is its own Fourier expansion.

THEOREM 26. (Uniqueness Theorem.) *Almost periodic functions which have the same Fourier expansion are equal.*

If G and H have the same expansion, then the expansion of $F=G-H$ vanishes identically. From

$$(25) \quad D_{\rho\sigma} \times F = M_y(D_{\rho\sigma}(xy^{-1})F(y)) = \sum_{\tau=1}^s D_{\rho\tau}(x) M_y(D_{\tau\sigma}(y^{-1})F(y))$$

it follows that $D_{\rho\sigma} \times F$ vanishes for any (D, ρ, σ) . Hence $\phi \times F = 0$ for any special weight function. But F is the limit of such functions (Theorem 23). Therefore $F = 0$.

PART III. THEOREMS ON APPROXIMATION AND SUMMATION

THEOREM 27. (Approximation Theorem.) *If F is almost periodic, it is the limit of a sequence of finite linear aggregates of the form $\sum f D_{\rho\sigma}$ with $f \in L$. More precisely, if D^N , $N = 1, 2, 3, \dots$, are the representations occurring in the Fourier expansion of F , if the s^N are their respective degrees, there exist elements $f_{N,\rho\sigma}^m$ of L ($m = 1, 2, 3, \dots$) such that for each m only a finite number of them is $\neq 0$, and that F is the limit, for $m \rightarrow \infty$, of the finite aggregates*

$$F_m = \sum_{N=1}^{\infty} \left(s^N \sum_{\rho, \sigma=1}^{s^N} f_{N,\rho\sigma}^m D_{\rho\sigma}^N \right).$$

F is the limit of a sequence $F_m = \phi_m \times F$, each ϕ_m being a special weight function. Using (25) we find that this sequence has the property stated in the theorem.

THEOREM 28. *Let the sequence F_m , $m = 1, 2, 3, \dots$, be part of a totally bounded or compact set of Ap . In order that the sequence have a limit it is sufficient that it be formally convergent.*

As the closure of a totally bounded set is compact (cf. [4], Theorem 11 and Definition 10), it is sufficient to consider the second case. Owing to the compactness, the sequence has at least one condensation point, and, using the compactness again, we have to show that it has no more than one condensation point. Otherwise there would exist two subsequences F_{p_m}, F_{q_m} of F_m , having two different limits, G and H respectively. The Fourier expansion of G is the formal limit of the Fourier expansions of the sequence F_{p_m} , and therefore, the sequence F_m being formally convergent, of the sequence F_m . Similarly, the Fourier expansion of H is the formal limit of the expansions of the sequence F_m . Thus G and H have the same Fourier expansion, and by Theorem 26, $G = H$, which contradicts our hypothesis.

DEFINITION 10. *The function $F \in Ap$ will be called a class function if for all x and y , $F(yxy^{-1}) = F(x)$, or, which is equivalent, if for all x and y , $F(yx) = F(xy)$.*

A formal series (24) will be called a class series if

$$(26) \quad f_{N,\rho\sigma} = f_N \delta_{\rho\sigma}.$$

THEOREM 29. *If $F(x) \in Ap$, then $F_0(x) = M_y F(yxy^{-1})$ is a class function, and if (24) is the Fourier expansion of $F(x)$, then the Fourier expansion of $F_0(x)$ is*

$$\sum_{N=1}^{\infty} s^N f_N \sum_{\rho=1}^s D_{\rho\rho}^N,$$

where

$$f_N = \frac{1}{s^N} \sum_{\rho=1}^s f_{N,\rho\rho}.$$

In order that a function F of Ap be a class function, it is necessary and sufficient that its Fourier expansion be a class series.

We have

$$F_0(xzx^{-1}) = M_y F(yzxz^{-1}y^{-1}) = M_y F(yxy^{-1}) = F_0(x);$$

hence $F_0(x)$ is a class function. For a fixed representation $\{D_{\rho\sigma}\}$ we have

$$\begin{aligned} M_x(D_{\rho\sigma}(x^{-1})F_0(x)) &= M_x M_y(D_{\rho\sigma}(x^{-1})F(yxy^{-1})) = M_x M_y(D_{\rho\sigma}(y^{-1}x^{-1}y)F(x)) \\ &= M_y \sum_{p,q=1}^s D_{\rho p}(y^{-1})D_{q\sigma}(y) \cdot M_x(D_{pq}(x^{-1})F(x)) \\ &= \frac{\delta_{\rho\sigma}}{s} \sum_{p=1}^s M_x(D_{pp}(x^{-1})F(x)) \end{aligned}$$

and this proves the statement about the Fourier series of $F_0(x)$.

If the Fourier expansion of F is a class series, the functions F and F_0 have the same Fourier expansion. By Theorem 26, $F = F_0$, and hence F is a class function, because F_0 is one.

Conversely if $F(x) = F(yxy^{-1})$, then $F(x) = M_y F(yxy^{-1}) = F_0(x)$, and by the first part of our theorem, the Fourier expansion of $F_0(x)$ is a class series.

THEOREM 30. (Summation Theorem.) *Let D^N , $N=1, 2, 3, \dots$, denote a sequence of irreducible normal representations, and the s^N their respective degrees. There exists a sequence of special weight functions ϕ_m , $m=1, 2, 3, \dots$, with the following properties:*

1. Each ϕ_m is a class function.
2. All Fourier coefficients of ϕ_m are ≥ 0 , ≤ 1 .
3. Any almost periodic function F whose Fourier expansion contains no other representations than the given ones, is the limit of the sequence $\phi_m \times F$, $m=1, 2, \dots$.

In particular, there exists a square array of numbers r_N^m , $m, N=1, 2, \dots$, with the following properties:

α . For each m only a finite number of them is $\neq 0$.

β . $0 \leq r_N^m \leq 1$.

γ . If an almost periodic function F has a Fourier expansion

$$\sum_{N=1}^{\infty} s^N \sum_{\rho, \sigma=1}^{s^N} f_{N, \rho \sigma} D_{\rho \sigma}^N$$

[any number of the coefficients $f_{N, \rho \sigma}$ may vanish], then F is the limit, $m \rightarrow \infty$, of the finite aggregates

$$F_m = \sum_{N=1}^{\infty} r_N^m s^N \sum_{\rho, \sigma=1}^{s^N} f_{N, \rho \sigma} D_{\rho \sigma}^N.$$

We determine numbers ϵ_N , all $\neq 0$, such that the series

$$(27) \quad \sum_{N=1}^{\infty} \epsilon_N s^N \left(\sum_{r=1}^{s^N} D_{rr}^N(x) \right)$$

is uniformly convergent, thus representing a numerical almost periodic function $f(x)$. There exist special weight functions $\chi_m(x)$, $m=1, 2, \dots$, such that the sequence of functions $f_m(x) = \chi_m \times f(x)$ is uniformly convergent to $f(x)$. By [1], Theorem 21,

$$f \times D_{\rho \sigma}^N = \epsilon_N D_{\rho \sigma}^N.$$

But $f_m \times D_{\rho \sigma}^N = \chi_m \times (f \times D_{\rho \sigma}^N)$. Considering $\epsilon_N \neq 0$, we conclude that

$$(28) \quad \lim_{m \rightarrow \infty} \chi_m \times D_{\rho \sigma}^N = D_{\rho \sigma}^N.$$

In particular,

$$\lim_{m \rightarrow \infty} M_x(D_{\rho \sigma}^N(x^{-1})\chi_m(x)) = \delta_{\rho \sigma}.$$

We now consider the class function

$$(29) \quad \psi_m(x) = M_y \chi_m(yxy^{-1});$$

$\psi_m(x)$ is obviously a weight function, and Theorem 29 shows that it has the form

$$(30) \quad \psi_m(x) = \sum_D s_D a_D^m \left(\sum_{r=1}^{s_D} D_{rr} \right)$$

with a finite number of terms, so that it is a special weight function. From (28) it follows that

$$\lim_{m \rightarrow \infty} a_D^m = 1 \quad \text{for } D = D^N.$$

The function

$$\psi_m' = \bar{\psi}_m(x^{-1}) = \sum_D s_D \bar{a}_D^m \left(\sum_{\tau=1}^{s_D} D_{\tau\tau} \right)$$

has all the properties of ψ_m . Finally we introduce the class function $\phi_m = \psi_m \times \psi_m'$, for which we find

$$\phi_m = \sum_D s_D r_D^m \left(\sum_{\tau=1}^{s_D} D_{\tau\tau} \right),$$

where

$$r_D^m = |a_D^m|^2,$$

with

$$(31) \quad \lim_{m \rightarrow \infty} r_D^m = 1 \quad \text{for } D = D^N.$$

It is easy to find that ϕ_m is again a weight function. Its coefficients r_D^m are ≥ 0 , and as r_D^m , $\phi_m(x)$ are real and ≥ 0 ,

$$r_D^m = M_x(D_{11}(x^{-1})\phi_m(x)) \leq M_x(|D_{11}(x^{-1})| \cdot |\phi_m(x)|) \leq M_x(\phi_m(x)) = 1.$$

Hence the properties 1 and 2 of our theorem are fulfilled. Using (31) we conclude from Theorem 24, 4, that for any F whose Fourier expansion contains no other representations than the given ones, the sequence $\phi_m \times F$ converges formally to F . By Theorem 21, this sequence is part of the compact set $((\mathfrak{R}_F)_{\text{conv}})_{\text{cl}}$. Thus we may apply Theorem 28, and this proves property 3.

The second half of Theorem 30 is an immediate consequence of the first half.

DEFINITION 11. *A system of irreducible normal representations will be called a module $\mathfrak{M}$ if it contains with every representation its complex conjugate, and the Fourier series of the product of any two occurring representation coefficients contains no other representations than the given ones.* A module will be called countable if it contains only a finite or enumerable number of representations.*

* If D^M, D^N are any two normal representations, then, in the sense in which these symbols are used in the theory of group-representations, the direct product $D^M D^N$ is a finite sum

$$\sum_P c_P^{MN} D^P$$

(the c_P^{MN} are the "composition coefficients"). Hence the product of any two representation coefficients has a finite Fourier expansion.

LEMMA 3. *There is always a smallest module containing a given set of representations; if the given set is countable, then so is also the smallest module.*

Consider all finite subsystems of the given set of representation coefficients and their complex conjugates and form all possible power-products with them. The totality of the representations occurring in their Fourier series is the desired module.

DEFINITION 12. *Given any module $\mathfrak{M}$, we shall denote by $\{\mathfrak{M}\}$ the set of Ap consisting of those functions whose Fourier expansion contains no other representations than those occurring in $\mathfrak{M}$.*

THEOREM 31. *Given $\mathfrak{M}$, the set $\{\mathfrak{M}\}$ has the following properties:*

1. *If $F(x)$ is in $\{\mathfrak{M}\}$, every $F(xa)$ is in $\{\mathfrak{M}\}$.*
2. *If $F(x)$ is in $\{\mathfrak{M}\}$, every $F(ax)$ is in $\{\mathfrak{M}\}$.*
3. *If $F(x)$ is in $\{\mathfrak{M}\}$, every $\alpha F(x)$ is in $\{\mathfrak{M}\}$.*
4. *If $F(x)$ and $G(x)$ are in $\{\mathfrak{M}\}$, $F(x) \pm G(x)$ is in $\{\mathfrak{M}\}$.*
5. *If $F_1, F_2, \dots$ are in $\{\mathfrak{M}\}$, and if F is the limit of F_n , then F is in $\{\mathfrak{M}\}$.*
6. *If $f_1, \dots, f_n$ are numerical functions of $\{\mathfrak{M}\}$, and $F(u_1, \dots, u_n)$ is a numerical function which is defined and uniformly continuous for the range of $f_1, \dots, f_n$, then $F(f_1(x), \dots, f_n(x))$ is also contained in $\{\mathfrak{M}\}$.*
7. *If at least one of the two functions $\alpha(x), F(x)$ is in $\{\mathfrak{M}\}$, then $\alpha \times F(x)$ is also. ($\alpha(x)$ is numerical.)*

1-5 follow easily from the formal properties of Fourier expansion. 7 is an immediate consequence of Theorem 24, 4. As regards 6, if f and g are numerical finite aggregates of representations from $\{\mathfrak{M}\}$, then $\bar{f}$ as well as fg is also $\epsilon\{\mathfrak{M}\}$ by definition of $\{\mathfrak{M}\}$. Applying the approximation theorem and 5, the same holds for $\bar{f}$ and fg for any numerical functions f, g from $\{\mathfrak{M}\}$. Hence 6 is true if $F(u_1, \dots, u_n)$ is a polynomial in $u_1, \dots, u_n$ and their complex conjugates. $F(u_1, \dots, u_n)$ being uniformly continuous on the range of $f_1, \dots, f_n$, which is a bounded set since $f_1, \dots, f_n$ are a.p., $F(f_1(x), \dots, f_n(x))$ is a uniform limit of polynomials in $f_1(x), \dots, f_n(x)$ and their complex conjugates. Thus 5 completes the proof of 6.

THEOREM 32. *If $\mathfrak{M} = (D^1, D^2, \dots)$ is a countable module, there exists in $\{\mathfrak{M}\}$ a sequence of special weight functions $\phi_1, \phi_2, \dots$ such that*

$$(i) \quad \phi_m = \sum_{N=1}^{\infty} s^N r_N^m \left(\sum_{\tau=1}^{s^N} D_{\tau\tau}^N \right);$$

$$(ii) \quad 0 \leq r_N^m \leq 1, \quad \lim_{m \rightarrow \infty} r_N^m = 1.$$

As in the proof of Theorem 30, we choose the $\epsilon_N \neq 0$ such that the series (27) is uniformly convergent, and construct the special weight functions $\phi_1, \phi_2, \dots$ mentioned therein.

By Theorem 30 these $\phi_1, \phi_2, \dots$ have the properties (i), (ii) of our theorem, so we need prove only that they belong to $\{\mathfrak{M}\}$.

As $f(x)$ belongs to $\{\mathfrak{M}\}$, for every finite subset $a_1, \dots, a_n$ of $\mathfrak{G}$, all $f(xa_1), \dots, f(xa_n)$ also do, and with them their continuous function

$$\begin{aligned} \max (|f(xa_1) - f(a_1)|, \dots, |f(xa_n) - f(a_n)|) \\ = \text{l.u.b.}_{y \in \mathfrak{G}} (|f(xy) - f(y)|). \end{aligned}$$

Since $f(z)$ is almost periodic, we can select a sequence of such sets $a_1, \dots, a_n$, so that this converges uniformly to the translation function

$$t(x) = \text{l.u.b.}_{y \in \mathfrak{G}} (|f(xy) - f(y)|).$$

Thus $t(x)$, as well as the $\psi_*(x) = \omega_*(t(x))$ of Theorem 22, which is a continuous function of $t(x)$, belongs to $\{\mathfrak{M}\}$, and with it

$$\psi(x) = \frac{\psi_*(x)}{M_x \psi_*(x)}.$$

The $\chi_1(x)$ of Lemma 2 contains only such representations as occur in the Fourier series of $\psi(x)$; therefore it also belongs to $\{\mathfrak{M}\}$, and with it $\chi(x)$. The same is true of the χ of Theorem 22 and also of the $\chi_1, \chi_2, \dots$ in the proof of Theorem 30 as well as of the $\psi_1, \psi_2, \dots$, since these contain only such representations as occur in the Fourier series of the corresponding functions $\chi_1, \chi_2, \dots$. It follows finally that $\phi_1, \phi_2, \dots \subset \{\mathfrak{M}\}$, and the proof is complete.

THEOREM 33. (Isolation Theorem.) *Let F be an almost periodic function, with a Fourier expansion*

$$(32) \quad \sum_D \left(s_D \sum_{\rho, \sigma=1}^{s_D} f_{D\rho\sigma} D_{\rho\sigma} \right),$$

and let $\mathfrak{M}$ be any module. There exists an almost periodic function, we shall denote it by $F_{\mathfrak{M}}$, whose Fourier expansion consists of exactly those terms

$$(33) \quad s_D \sum_{\rho, \sigma=1}^{s_D} f_{D\rho\sigma} D_{\rho\sigma}$$

of (32), for which D is contained in $\mathfrak{M}$. And we have $F_{\mathfrak{M}} \rightarrow F$.

The given module $\mathfrak{M}$ need not be enumerable from the outset, but we may

replace it by the enumerable module generated by those representations of $\mathfrak{M}$ which occur in (32). Hence $\mathfrak{M}$ may be assumed to be enumerable. Let it contain the representations $D^1, D^2, \dots$. With the special weight functions $\phi_1, \phi_2, \dots$ from Theorem 32 we construct the functions $\phi_1 \times F, \phi_2 \times F, \dots$. By Theorem 21 the functions $\phi_m \times F$ are all $\sim F$, and by the properties of ϕ_m , the functions $\phi_m \times F$ are formally convergent. By Theorem 28 the sequence $\phi_m \times F$ has a limit function $f_{\mathfrak{M}}$, and by Theorems 24, 4 and 32 the Fourier series of $f_{\mathfrak{M}}$ has the stated form.

THEOREM 34. *Let F be an almost periodic function and $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ a sequence of monotonically increasing modules which in their sum contain all representations occurring in F .*

The sequence of functions

$$(34) \quad F_{\mathfrak{M}_n} \quad (n = 1, 2, \dots)$$

converges to F .

The functions (34) converge formally to F , and are all $\sim F$. By Theorem 28, F is the limit of the sequence.

THEOREM 35. (Decomposition Theorem.) *Let F be an almost periodic function and let it be possible to divide the representations occurring in F in a sequence of systems $\mathfrak{s}_1, \mathfrak{s}_2, \mathfrak{s}_3, \dots$, in such a way that for each $k (= 1, 2, 3, \dots)$, the least module containing $\mathfrak{S}(\mathfrak{s}_1, \mathfrak{s}_2, \dots, \mathfrak{s}_k)$ has no element in common with $\mathfrak{S}(\mathfrak{s}_{k+1}, \mathfrak{s}_{k+2}, \dots)$. There exists, for each k , an almost periodic function, we shall denote it by $F_{\mathfrak{s}_k}$, whose Fourier expansion consists of exactly those terms (33) of (32), for which D is contained in $\mathfrak{s}_k$; and the series*

$$F_{\mathfrak{s}_1} + F_{\mathfrak{s}_2} + \dots$$

converges to F .

If we denote the smallest module containing $\mathfrak{s}_1, \dots, \mathfrak{s}_k$ by $\mathfrak{M}_k$, the desired function $F_{\mathfrak{s}_k}$ is $F_{\mathfrak{M}_k}$ for $k=1$, and $F_{\mathfrak{M}_k} - F_{\mathfrak{M}_{k-1}}$ for $k \geq 2$, and our theorem reduces to Theorem 34.

PART IV. APPLICATIONS TO HILBERT SPACE

In this section we shall assume L to be either a Hilbert space $\mathfrak{H}$, or the space $\mathfrak{B}$ of all bounded transformations in $\mathfrak{H}$. We shall consider these spaces with the help of the various topologies which have been discussed in [4], Chapter IV, as well as in the places quoted there. According to whether we consider $\mathfrak{H}$ in the topology based on the neighborhoods $\mathfrak{U}_1$ or $\mathfrak{U}_2$, we shall denote $\mathfrak{H}$ by $\mathfrak{H}_1$ or $\mathfrak{H}_2$ respectively. As for $\mathfrak{B}$, corresponding to the neighborhoods $\mathfrak{U}_3, \mathfrak{U}_4, \mathfrak{U}_5$ we shall denote it by $\mathfrak{B}_3, \mathfrak{B}_4, \mathfrak{B}_5$, respectively.

THEOREM 36. (α) If $F(x) \in \mathfrak{S}^{\mathfrak{G}}$ is almost periodic in $\mathfrak{S}_1^{\mathfrak{G}}$ it is also almost periodic in $\mathfrak{S}_2^{\mathfrak{G}}$. (β) If $F(x) \in \mathfrak{B}^{\mathfrak{G}}$ is almost periodic in $\mathfrak{B}_3^{\mathfrak{G}}$ it is also almost periodic in $\mathfrak{B}_4^{\mathfrak{G}}$, and (γ) if it is almost periodic in $\mathfrak{B}_4^{\mathfrak{G}}$ it is also almost periodic in $\mathfrak{B}_5^{\mathfrak{G}}$.

(δ) $F(x)$ is almost periodic in $\mathfrak{S}_2^{\mathfrak{G}}$ if and only if the numerical function $(F(x), g)$, which is the inner product of $F(x)$ with the constant $g \in \mathfrak{S}$, is almost periodic for every $g \in \mathfrak{S}$.

(ε) $F(x)$ is almost periodic in $\mathfrak{B}_4^{\mathfrak{G}}$ if and only if $F(x)f$, that is, the value of $F(x)$ operated upon the constant $f \in \mathfrak{S}$, is almost periodic in $\mathfrak{S}_1^{\mathfrak{G}}$ for every $f \in \mathfrak{S}$.

(ζ) $F(x)$ is almost periodic in $\mathfrak{B}_5^{\mathfrak{G}}$ if and only if $F(x)f$ is almost periodic in $\mathfrak{S}_2^{\mathfrak{G}}$ for every $f \in \mathfrak{S}$, that is, if and only if the numerical function $(F(x)f, g)$ is almost periodic for every pair $f, g, \in \mathfrak{S}$.

Ad (α). The topology of $\mathfrak{S}_2$ is weaker than the topology of $\mathfrak{S}_1$. More precisely: to any $U_2 \in \mathfrak{U}_2$ there corresponds a $U_1 \in \mathfrak{U}_1$ such that $U_1 \subset U_2$. Hence, if a set $S \subset \mathfrak{S}$ is totally bounded in $\mathfrak{S}_1$ it is also totally bounded in $\mathfrak{S}_2$. And this proves the proposition, if we remember Definition 1.

Ad (β) and (γ). A similar argument as ad (α).

Ad (δ). Let S be any set of $\mathfrak{S}_2^{\mathfrak{G}}$. If $g \in \mathfrak{S}$ we denote by S_g the set of numerical functions $(F(x), g)$, $F \in S$, and for any finite number of elements $g_1, \dots, g_n \in \mathfrak{S}$ we denote by $S_{g_1, \dots, g_n}$ the set of n -dimensional vector functions with components $(F(x), g_1), \dots, (F(x), g_n)$, $F \in S$. It is easy to see that S is totally bounded in $\mathfrak{S}_2^{\mathfrak{G}}$ if and only if all sets $S_{g_1, \dots, g_n}$ are totally bounded. On the other hand, we have to prove that S is totally bounded in $\mathfrak{S}_2$ if and only if all S_g are totally bounded. Thus it is sufficient to prove that, for a fixed S , the total boundedness of all sets S_g implies the total boundedness of all sets $S_{g_1, \dots, g_n}$. This follows from [4], Theorem 9, if we observe that a vector $\{\phi_1, \dots, \phi_n\}$ may be considered as a continuous function of its components $\phi_1, \dots, \phi_n$ in the obvious topology of vector spaces.

Ad (ε) and (ζ). A similar argument as ad (δ).

THEOREM 37. Let $F(x)$ be an almost periodic function, its Fourier expansion being (throughout this part, we will write the Fourier coefficients $\alpha_{\rho\sigma}(D)$ without the factors s_D , the degree of D)

$$(35) \quad F(x) \sim \sum_{D, \rho, \sigma} \alpha_{\rho\sigma}(D) D_{\rho\sigma}(x).$$

If $F(x)$ is in $\mathfrak{S}^{\mathfrak{G}}$, we have $\alpha_{\rho\sigma}(D) \in \mathfrak{S}$, and $(F(x), g)$, for any $g \in \mathfrak{S}$, has the Fourier expansion

$$(36) \quad (F(x), g) \sim \sum_{D, \rho, \sigma} (\alpha_{\rho\sigma}(D), g) D_{\rho\sigma}(x).$$

If $F(x)$ is in $\mathfrak{B}^{\mathfrak{G}}$, we have $\alpha_{\rho\sigma}(D) \in \mathfrak{B}$; for any $f \in \mathfrak{G}$, $F(x)f$ has the Fourier expansion

$$(37) \quad F(x)f \sim \sum_{D, \rho, \sigma} (\alpha_{\rho\sigma}(D)f) \cdot D_{\rho\sigma}(x),$$

and for any $f, g \in \mathfrak{G}$, $(F(x)f, g)$ has the Fourier expansion

$$(38) \quad (F(x)f, g) \sim \sum_{D, \rho, \sigma} (\alpha_{\rho\sigma}(D)f, g) \cdot D_{\rho\sigma}(x).$$

Remembering the definition of $\alpha_{\rho\sigma}(D)$, it is sufficient to prove

$$(39) \quad M_z(F(x), g) = (M_z F(x), g)$$

in the case of (36), and

$$(40) \quad M_z(F(x)f) = (M_z F(x))f$$

in the case (37); (38) follows from their combination.

The mean with respect to the strong topology ($\mathfrak{G}_1$ or $\mathfrak{B}_4$) has, a fortiori, all the properties of the mean in the weak topology ($\mathfrak{G}_2$ or $\mathfrak{B}_5$). Hence, by the uniqueness property of the mean (Theorem 17), it is sufficient to consider the cases of weak topology. In these cases the relations (39), (40) follow by a new application of Theorem 17, since $(M_z F(x), g)$ and $(M_z F(x))f$ have the properties required in this theorem.

THEOREM 38. *If $F(x)$ is an almost periodic function of the class $\mathfrak{B}_\nu^{\mathfrak{G}}$ ($\nu = 3, 4, 5$), and $\alpha \in \mathfrak{B}$, then $\alpha F(x)$ and $F(x)\alpha$ [αF and $F\alpha$ are the operational products of F and α] are again almost periodic functions of the same class; and the Fourier expansions of αF and $F\alpha$ may be obtained from the Fourier expansion of $F(x)$ by term-by-term multiplication with α on the right or on the left.*

The first half of the theorem follows from the fact that, for a fixed α , αF and $F\alpha$ are continuous functions of F in each of the topologies $\mathfrak{B}_3, \mathfrak{B}_4, \mathfrak{B}_5$. The second half (in the case of αF) follows from Theorem 37, Formula (38), if, in the formula of our statement connecting the expansions of αF and F , we apply (38) and then replace F, f, g by $\alpha F, f, g$ in the left-hand member, and by $F, f, \alpha^* g$ (α^* is the adjoint of α) in the right-hand member and compare the results; in the case of $F\alpha$ we replace F, f, g by $F\alpha, f, g$ and by $F, \alpha f, g$ respectively.

THEOREM 39. Let $F(x)$ be an almost periodic function of $\mathfrak{B}_1^{\mathfrak{G}}$; let (35) be its Fourier expansion. If $F(x)$ is a representation of the group by means of unitary transformations, that is, if $F(xy) = F(x)F(y)$, $F(1) = 1$, $F(x^{-1}) = F(x)^*$, then we have the following:

(i) $F(x)$ is almost periodic as a function of $\mathfrak{B}_1^{\mathfrak{G}}$; moreover, F is the limit, in the $\mathfrak{B}_1$ -topology, of the sequence

$$(41) \quad F_m = \sum_{N=1}^m \left(\sum_{\rho, \sigma=1}^{s_D N} \alpha_{\rho\sigma}(D^N) D_{\rho\sigma}^N(x) \right).$$

(ii) The Fourier coefficients have the properties

$$(42) \quad \alpha_{\rho\sigma}(D) \alpha_{\tau\nu}(E) = 0 \quad \text{if } D \neq E,$$

$$(43) \quad \alpha_{\rho\sigma}(D) \alpha_{\tau\nu}(D) = \delta_{\sigma\tau} \alpha_{\rho\nu}(D) \quad [\rho, \sigma, \tau, \nu = 1, \dots, s_D],$$

$$(44) \quad \alpha_{\rho\sigma}(D)^* = \alpha_{\sigma\rho}(D) \quad [\rho, \sigma = 1, \dots, s_D].$$

(iii) The system of operators $\alpha_{\rho\rho}(D)$ is a resolution of the identity, that is to say, each of them is a projection operator, any two of them are orthogonal, and their sum is the identity.

Conversely, if a set of elements $\alpha_{\rho\sigma}(D)$ of $\mathfrak{B}$ fulfills the conditions (ii) and (iii), then the series (35) is the Fourier expansion of an almost periodic function $F(x)$ with the assigned properties.

Remark concerning the theorem. Before proving the theorem we want to point out the algebraic aspect of the proposition.

Since $\alpha_{\rho\rho}(D)$ is a projection it corresponds to a closed linear manifold, say $\mathfrak{M}_\rho^D$. Select an orthonormal system determining $\mathfrak{M}_\rho^D$: $\psi_{11}^D, \psi_{12}^D, \dots$ (the sequence is empty, finite, or countable). Now, (43) and (44) imply

$$\alpha_{\rho 1}(D)^* \alpha_{\rho 1}(D) = \alpha_{11}(D), \quad \alpha_{\rho 1}(D) \alpha_{\rho 1}(D)^* = \alpha_{\rho\rho}(D),$$

and this means that $\alpha_{\rho 1}(D)$ maps $\mathfrak{M}_1^D$ in a one-to-one and unitary way on $\mathfrak{M}_\rho^D$, while $\alpha_{\rho 1}(D)f \equiv 0$ if f is orthogonal to $\mathfrak{M}_1^D$. Thus, if we define $\psi_{\rho\nu}(D) = \alpha_{\rho 1}(D)\psi_{1\nu}(D)$, the $\psi_{\rho 1}(D), \psi_{\rho 2}(D), \dots$ determine $\mathfrak{M}_\rho^D$. By (43) we have $\alpha_{\rho\rho}(D)\alpha_{\tau 1}(D) = \alpha_{\rho 1}(D)$, and this implies $\alpha_{\rho\sigma}(D)\psi_{\sigma\nu}(D) = \psi_{\rho\nu}(D)$.

Since, by (iii), $\sum_{D, \rho} \alpha_{\rho\rho}(D) = 1$, the $\mathfrak{M}_\rho^D$ together determine $\mathfrak{G}$, and, since, by (42), (43), (44), $\alpha_{\sigma\rho}(D)^* \alpha_{\tau\nu}(E) = 0$ if $D \neq E$ or $\sigma \neq \tau$, the $\mathfrak{M}_\rho^D$ are mutually orthogonal. Thus the $\psi_{\rho\nu}(D)$ form a complete orthonormal system in $\mathfrak{G}$. For this system we found

$$\alpha_{\tau\sigma}(E) \psi_{\rho\nu}(D) = \begin{cases} 0, & \text{if } D \neq E \text{ or } \nu \neq \rho, \\ \psi_{\tau\nu}(D) & \text{if } D = E, \nu = \rho. \end{cases}$$

Remembering that the $\alpha_{\tau\nu}(D)$ form the Fourier expansion of $F(x)$, we obtain the formulas

$$F(x)\psi_{\rho\nu}(D) = \sum_{\tau=1}^{s_D} D_{\tau\rho}(x)\psi_{\tau\nu}(D).$$

Thus our theorem proves that the representation $F(x)$ may be reduced to split up according to the finite irreducible representations $D(x)$. The subspaces which correspond to this reduction are determined by the

$$(45) \quad \psi_{1\nu}(D), \psi_{2\nu}(D), \dots, \psi_{s_D\nu}(D) \text{ for all } D, \nu. —$$

We pass on to the proof. Let $F(x)$ have the assumed properties. $F(xy)$ is an almost periodic function in (x, y) and for each y an almost periodic function in x . From the approximating properties of the Fourier expansion it follows that the Fourier expansion of $F_y(x) = F(xy)$ as x -function may be derived from the series (35) in a formal way, namely

$$(46) \quad F_y(x) \sim \sum_{D, \rho, \sigma} \alpha_{\rho\sigma}(D; y) D_{\rho\sigma}(x),$$

where

$$(47) \quad \alpha_{\rho\sigma}(D; y) = \sum_{\tau=1}^{s_D} \alpha_{\rho\tau}(D) D_{\sigma\tau}(y).$$

On the other hand, $F_y(x) = F(x)F(y)$, and by Theorem 38,

$$(48) \quad F_y(x) \sim \sum \alpha_{\rho\sigma}(D) F(y) D_{\rho\sigma}(x).$$

A comparison of (46) and (48) yields, for any D, ρ, σ ,

$$(49) \quad \alpha_{\rho\sigma}(D) F(y) = \sum_{\tau=1}^{s_D} \alpha_{\rho\tau}(D) D_{\sigma\tau}(y).$$

Now we take y variable, and we again apply Theorem 38. This gives the relations (42) and (43). In order to prove (44) we need only compare the relations

$$(50) \quad F(x^{-1}) \sim \sum \alpha_{\rho\sigma}(D) D_{\rho\sigma}(x^{-1}) = \sum \alpha_{\rho\sigma}(D) \overline{D_{\sigma\rho}(x)},$$

$$(51) \quad F(x)^* \sim \sum \alpha_{\rho\sigma}(D)^* \overline{D_{\rho\sigma}(x)}$$

(which follow from the approximating properties of the Fourier series), and observe that the representations $\{\overline{D_{\rho\sigma}}\}$ also form a set of irreducible normal representations of $\mathfrak{G}$.

If we apply (44) to $\rho = \sigma$ we find that $\alpha_{\rho\rho}(D)$ is self-adjoint, and (43) gives $(\alpha_{\rho\rho}(D))^2 = \alpha_{\rho\rho}(D)$. Hence, each $\alpha_{\rho\rho}(D)$ is a projection. It follows from (42) and (43) that any two of them are orthogonal. Hence any finite number of

them, and also their (infinite) sum, is a projection again. We still have to prove that $\sum_{D, \rho} \alpha_{\rho\rho}(D)$ is the unity and that (i) holds.

Let f be any element of $\mathfrak{F}$. We want to evaluate the difference

$$F_p(x)f - F_q(x)f, \quad p > q,$$

$F_p(x)$ being given by (41). Let D be any representation occurring in $F_p(x)f - F_q(x)f$. Writing

$$\beta^D(x) = \sum_{\rho, \sigma=1}^{s_D} \alpha_{\rho\sigma}(D) D_{\rho\sigma}(x),$$

we have

$$\|\beta(x)f\|^2 = (\beta(x)f, \beta(x)f) = \sum_{\rho, \sigma, \tau, \nu=1}^{s_D} D_{\rho\sigma}(x) \overline{D_{\tau\nu}(x)} (\alpha_{\rho\sigma}(D)f, \alpha_{\tau\nu}(D)f);$$

but

$$(\alpha_{\rho\sigma}(D)f, \alpha_{\tau\nu}(D)f) = (\alpha_{\tau\nu}(D)^* \alpha_{\rho\sigma}(D)f, f) = (\alpha_{\nu\tau}(D) \alpha_{\rho\sigma}(D)f, f) = \delta_{\rho\sigma} (\alpha_{\nu\sigma}(D)f, f);$$

hence

$$\begin{aligned} \|\beta(x)f\|^2 &= \sum_{\nu, \sigma=1}^{s_D} \left(\sum_{\tau=1}^{s_D} D_{\tau\sigma}(x) \overline{D_{\tau\nu}(x)} \right) (\alpha_{\nu\sigma}(D)f, f) \\ &= \sum_{\nu, \sigma=1}^{s_D} \delta_{\nu\sigma} (\alpha_{\nu\sigma}(D)f, f) \\ &= \sum_{\rho=1}^{s_D} (\alpha_{\rho\rho}(D)f, f) = \sum_{\rho=1}^{s_D} \|\alpha_{\rho\rho}(D)f\|^2. \end{aligned}$$

Thus

$$(52) \quad \|F_p(x)f - F_q(x)f\|^2 = \sum_{N=q+1}^p \sum_{\rho=1}^{s_{D^N}} \|\alpha_{\rho\rho}(D^N)f\|^2.$$

Since the $\alpha_{\rho\rho}(D)$ are mutually orthogonal projections and the right side of (52) is independent of x , it follows easily that the sequence of functions $F_m(x)$ is uniformly convergent in the topology $\mathfrak{B}_4$. Thus $F(x)$ is the limit of the sequence (41) and is almost periodic also in the class $\mathfrak{B}_4^{\mathfrak{G}}$. Now

$$F(1) = \lim_{m \rightarrow \infty} F_m(1) = \lim_{m \rightarrow \infty} \sum_{N=1}^m \sum_{\rho=1}^{s_{D^N}} \alpha_{\rho\rho}(D^N) = \sum_{D, \rho} \alpha_{\rho\rho}(D).$$

But $F(1) = 1$, by assumption, and this proves the last missing part of (iii).

Conversely, let a set of elements $\alpha_{\rho\sigma}(D)$ of $\mathfrak{B}$ satisfy (ii) and (iii). As before, we deduce the relation (52). Hence the sequence $F_m(x)$, defined by

(41), is convergent, in the topology $\mathfrak{B}_4$, to an almost periodic function $F(x)$ whose Fourier expansion is the uniformly convergent series (35). The group properties of $F(x)$ are easily verified from its series (35) on the basis of the known properties (ii), (iii).

THEOREM 40. *Let the group $\mathfrak{G}$ be a metric, locally compact, separable space in which the product xy is a continuous function of (x, y) . We consider in $\mathfrak{G}$ a right invariant Haar measure and the Hilbert space $\mathfrak{H}$ consisting of all Lebesgue measurable functions of integrable square in $\mathfrak{G}$, and let $\mathfrak{T}(x)$ denote, for each x , the unitary operator which transforms every element $f(t) \in \mathfrak{H}$ into $f(tx)$.*

In order that $\mathfrak{G}$ be compact it is necessary and sufficient that $\mathfrak{T}(x)$ be an almost periodic function of $\mathfrak{B}_5^$.*

Obviously $\mathfrak{T}(x)$ has the properties

$$\mathfrak{T}(x)^{-1} = \mathfrak{T}(x)^*, \quad \mathfrak{T}(xy) = \mathfrak{T}(x)\mathfrak{T}(y), \quad \mathfrak{T}(1) = 1.$$

It follows easily that $\mathfrak{T}(x)$ is continuous in the topology $\mathfrak{B}_5$.

If $\mathfrak{G}$ is compact, every continuous function is almost periodic and this proves the necessity of our condition. Conversely, let $\mathfrak{T}(x)$ be almost periodic. We consider its Fourier expansion and the complete orthonormal system $\psi_\rho(D)$ constructed in the remark to Theorem 39. Consider one of its non-empty subsystems (45) and denote it by $\psi_1, \dots, \psi_s$. By the last relation of the remark we have

$$(53) \quad F(x)\psi_\rho(t) = \psi_\rho(xt) = \sum_{\tau} D_{\tau\rho}(x)\psi_\tau(t),$$

except perhaps for a t -set of measure 0, depending on x . By the theorem of Fubini there is a value $t=t_0$ for which (53) holds for all x except a set of measure 0. As $D(x)$ is unitary we obtain

$$\sum_{\rho} |\psi_\rho(xt_0)|^2 = \sum_{\rho} |\psi_\rho(t_0)|^2 = C.$$

Upon integrating over $\mathfrak{G}$, the left side gives s_D whereas the right side is C times the total measure of $\mathfrak{G}$. Thus $s_D > 0$ implies that $C \neq 0$ and that the total measure of $\mathfrak{G}$ is finite. If the total measure of $\mathfrak{G}$ is finite there cannot exist an $\epsilon_0 > 0$ and an infinite number of points on $\mathfrak{G}$ any two of which have a distance $> \epsilon_0$. This implies that $\mathfrak{G}$ is totally bounded. Being locally compact, $\mathfrak{G}$ is also complete. Hence the compactness of $\mathfrak{G}$.

BIBLIOGRAPHY

- [1] J. von Neumann, *Almost periodic functions in a group*, I, these Transactions, vol. 36 (1934), pp. 445–492.
- [2] S. Bochner, *Abstrakte fastperiodische Funktionen*, Acta Mathematica, vol. 61 (1933), pp. 149–184.
- [3] S. Bochner, *Fastperiodische Lösungen der Wellengleichung*, Acta Mathematica, vol. 62 (1934), pp. 227–237.
- [4] J. von Neumann, *On complete topological spaces*, these Transactions, vol. 37 (1935), pp. 1–20.
- [5] S. Bochner, *Beiträge zur Theorie der fastperiodischen Funktionen*, I, Mathematische Annalen, vol. 96 (1926), pp. 119–147.

COMPARISON OF CELLS

Reviewed by G. Hunt

JOHN VON NEUMANN dealt with the equivalence of measure spaces in several contexts. The present manuscript proves in detail the following theorem, which is well known for linear intervals:

Theorem: Let P be an n -cell (that is, a compact n -dimensional parallelepiped) and let ρ, σ be Radon measures on P which have the same finite total mass, which attribute no mass to the boundary of P or to any set reduced to a single point, and which attribute strictly positive mass to every non-empty open set. Then, there is a homeomorphism ϕ of P on itself which leaves fixed each point of the boundary and which carries ρ into σ (in the sense that $\rho(\phi(M)) = \sigma(M)$ for every Borel set M).

For $n > 1$, a set whose dimension lies between 0 and n may carry part of the mass of ρ or σ . This fact renders the theorem implausible, at first glance, and causes all the complications of the proof.

BIBLIOGRAPHY OF JOHN VON NEUMANN

BOOKS

Mathematische Grundlagen der Quantenmechanik. Berlin, Springer (1932) ; New York, Dover Publications (1943) ; Presses Universitaires de France (1947) ; Madrid, Instituto de Matematicas "Jorge Juan" (1949) ; trans. from German by ROBERT T. BEYER, Princeton University Press (1955).

With O. MORGENSTERN. *Theory of Games and Economic Behavior.* Princeton University Press (1944, 1947, 1953). 625 pp. German trans. by DOCQUIER (in press).

Functional Operators. Vol. I : *Measures and Integrals* (*Ann. Math. Studies*, No. 21, 261 pp.) ; Vol. II : *The Geometry of Orthogonal Spaces* (*Ann. Math. Studies*, No. 22, 107 pp.). Princeton University Press (1950).

The Computer and the Brain (Silliman Lectures). Yale University Press (1958). 82 p.p

Continuous Geometry (with an introduction by I. HALPERIN). Princeton University Press (1960) 229 pp.

ARTICLES*

1922

1. With M. FEKETE. Über die Lage der Nullstellen gewisser Minimumpolynome. *Jahresb.* 31:125-138 [Vol. I, No. 2]

1923

2. Zur Einführung der transfiniten Zahlen. *Acta Szeged.* 1:199-208 [Vol. I, No. 3]

1925

3. Eine Axiomatisierung der Mengenlehre. *J. f. Math.* 154:219-240 [Vol. I, No. 4]
4. Egenletesen sürü számsorozatok. *Math. Phys. Lapok* 32:32-40 [Vol. I, No. 5]

1926

5. Zur Prüferschen Theorie der idealen Zahlen. *Acta Szeged.* 2:193-227 [Vol. I, No. 6]

1927

6. With D. HILBERT and L. NORDHEIM. Über die Grundlagen der Quantenmechanik. *Math. Ann.* 98:1-30 [Vol. I, No. 7]
7. Zur Theorie der Darstellungen kontinuierlicher Gruppen. *Sitzungsber. d. Preuss Akad.* pp. 76-90 [Vol. I, No. 8]
8. Mathematische Begründung der Quantenmechanik. *Gött. Nach.* pp. 1-57 [Vol. I, No. 9]
9. Wahrscheinlichkeitstheoretischer Aufbau der Quantenmechanik. *Gött. Nach.* pp. 245-272 [Vol. I, No. 10]
10. Thermodynamik quantenmechanischer Gesamtheiten. *Gött. Nach.* pp. 273-291 [Vol. I, No. 11]
11. Zur Hilbertschen Beweistheorie. *Math. Zschr.* 26:1-46 [Vol. I, No. 12]

*With most items listed there is given a volume and a number. These denote the volume and the order number in that volume where the item may be found. Thus, [Vol. I, No. 2] at the end of item 1 implies that item 1 is the second article in Volume I.

1928

12. Die Zerlegung eines Intervalles in abzählbar viele kongruente Teilmengen. *Fund. Math.* **11**:230-238 [Vol. I, No. 13]
13. Ein System algebraisch unabhängiger Zahlen. *Math. Ann.* **99**:134-141 [Vol. I, No. 14]
14. Über die Definition durch transfinite Induktion, und verwandte Fragen der allgemeinen Mengenlehre. *Math. Ann.* **99**:373-391 [Vol. I, No. 15]
15. † Zur Theorie der Gesellschaftsspiele. *Math. Ann.* **100**:295-320 [Vol. VI, No. 1]
16. Die Axiomatisierung der Mengenlehre. *Math. Zschr.* **27**:669-752 [Vol. I, No. 16]
17. With E. WIGNER. Zur Erklärung einiger Eigenschaften der Spektren aus der Quantenmechanik des Drehelektrons, I. *Zschr. f. Phys.* **47**:203-220 [Vol. I, No. 18]
18. Einige Bemerkungen zur Diracschen Theorie des Drehelektrons. *Zschr. f. Phys.* **48**:868-881 [Vol. I, No. 17]
19. With E. WIGNER. Zur Erklärung einiger Eigenschaften der Spektren aus der Quantenmechanik des Drehelektrons, II. *Zschr. f. Phys.* **49**:73-94 [Vol. I, No. 19]
20. With E. WIGNER. Zur Erklärung einiger Eigenschaften der Spektren aus der Quantenmechanik des Drehelektrons, III. *Zschr. f. Phys.* **51**:844-858 [Vol. I, No. 20]

1929

21. Über eine Widerspruchsfreiheitsfrage der axiomatischen Mengenlehre. *J. f. Math.* **160**:227-241 [Vol. I, No. 21]
22. Über die analytischen Eigenschaften von Gruppen linearer Transformationen und ihrer Darstellungen. *Math. Zschr.* **30**:3-42 [Vol. I, No. 22]
23. With E. WIGNER. Über merkwürdige diskrete Eigenwerte. *Phys. Zschr.* **30**:465-467 [Vol. I, No. 23]
24. With E. WIGNER. Über das Verhalten von Eigenwerten bei adiabatischen Prozessen. *Phys. Zschr.* **30**:467-470 [Vol. I, No. 24]
25. Beweis des Ergodensatzes und des H-Theorems in der neuen Mechanik. *Zschr. f. Phys.* **57**:30-70 [Vol. I, No. 25]
26. Zur allgemeinen Theorie des Masses. *Fund. Math.* **13**:73-116 [Vol. I, No. 26]
27. Zusatz zur Arbeit "Zur allgemeinen . . ." *Fund. Math.* **13**:333 [Vol. I, No. 27]
28. Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren. *Math. Ann.* **102**:49-131 [Vol. II, No. 1]
29. Zur algebra der Funktionaloperatoren und Theorie der normalen Operatoren. *Math. Ann.* **102**:370-427 [Vol. II, No. 2]
30. Zur Theorie der unbeschränkten Matrizen. *J. f. Math.* **161**:208-236 [Vol. II, No. 3]

1930

31. Über einen Hilfssatz der Variationsrechnung. *Hamb. Abh.* **8**:28-31 [Vol. II, No. 4]

1931

32. Über Funktionen von Funktionaloperatoren. *Ann. Math.* **32**:191-226 [Vol. II, No. 5]
33. Algebraische Repräsentanten der Funktionen "bis auf eine Menge vom Masse Null". *J. f. Math.* **165**:109-115 [Vol. II, No. 6]
34. Die Eindeutigkeit der Schrödingerschen Operatoren. *Math. Ann.* **104**:570-578 [Vol. II, No. 7]
35. Bemerkungen zu den Ausführungen von Herrn St. Lesniewski über meine Arbeit "Zur Hilbertschen Beweistheorie". *Fund. Math.* **17**:331-334 [Vol. II, No. 8]
36. Die formalistische Grundlegung der Mathematik. Report to the Congress, Königsberg, September, 1931. *Erkenntnis* **2**:116-121 [Vol. II, No. 9]

1932

37. Zum Beweise des Minkowskischen Satzes über Linearformen. *Math. Zschr.* **30**:1-2 [Vol. II, No. 10]
38. Über adjungierte Funktionaloperatoren. *Ann. Math.* **33**:294-310 [Vol. II, No. 11]
39. Proof of the Quasi-ergodic Hypothesis. *N.A.S. Proc.* **18**:70-82 [Vol. II, No. 12]
40. Physical Applications of the Ergodic Hypothesis. *N.A.S. Proc.* **18**:263-266 [Vol. II, No. 13]
41. With B. O. KOOPMAN. Dynamical Systems of Continuous Spectra. *N.A.S. Proc.* **18**:255-263 [Vol. II, No. 14]
42. Über einen Satz von Herrn M. H. Stone. *Ann. Math.* **33**:567-573 [Vol. II, No. 15]

†Translated by S. BARGMANN. On the Theory of Games of Strategy. In : *Contributions to the Theory of Games*, Vol. IV (*Ann. Math. Studies*, No. 40, Princeton University Press 1959), pp. 13-42.

43. Einige Sätze über messbare Abbildungen. *Ann. Math.* 33:574-586 [Vol. II, No. 16]
44. Zur Operatorenmethode in der klassischen Mechanik. *Ann. Math.* 33:587-642 [Vol. II, No. 17]
45. Zusätze zur Arbeit "Zur Operatorenmethode . . ." *Ann. Math.* 33:789-791. See also No. 44 [Vol. II, No. 18]

1933

46. Die Einführung analytischer Parameter in topologischen Gruppen. *Ann. Math.* 34:170-190 [Vol. II, No. 19]
47. A koordináta-mérés pontosságának határai az elektron Dirac-féle elméletében (Über die Grenzen der Koordinatenmessungs-Genauigkeit in der Diracschen Theorie des Elektrons). *Mat. és Természettud . . . Értesítő*, 50:366-385 [Vol. II, No. 20]

1934

48. With P. JORDAN and E. WIGNER. On an Algebraic Generalization of the Quantum Mechanical Formalism. *Ann. Math.* 35:29-64 [Vol. II, No. 21]
49. Zum Haarschen Mass in topologischen Gruppen. *Compos. Math.* 1:106-114. [Vol. II, No. 22]
50. Almost Periodic Functions in a Group. I. *Amer. Math. Soc.* 36:445-492 [Vol. II, No. 23]
51. With A. H. TAUB and O. VELEN. The Dirac Equation in Projective Relativity. *N.A.S. Proc.* 20:383-388 [Vol. II, No. 24]

1935

52. On Complete Topological Spaces. *Amer. Math. Soc. Trans.* 37:1-20 [Vol. II, No. 25]
53. With S. BOCHNER. Almost Periodic Functions in Groups. II. *Amer. Math. Soc. Trans.* 37:21-50 [Vol. II, No. 26]
54. With S. BOCHNER. On Compact Solutions of Operational-Differential Equations. I. *Ann. Math.* 36:255-291 [Vol. IV, No. 1]
55. Charakterisierung des Spektrums eines Integraloperators. *Actualités Scient. et Ind. Series*, No. 229. Exposés Math., publiés à la mémoire de J. Herbrand, No. 13. Paris. 20 pp. [Vol. IV, No. 2]
56. On Normal Operators. *N.A.S. Proc.* 21:366-369 [Vol. IV, No. 3]
57. With P. JORDAN. On Inner Products in Linear, Metric Spaces. *Ann. Math.* 36:719-723 [Vol. IV, No. 4]
58. With M. H. STONE. The Determination of Representative Elements in the Residual Classes of a Boolean Algebra. *Fund. Math.* 25:353-378 [Vol. IV, No. 5]

1936

59. On a Certain Topology for Rings of Operators. *Ann. Math.* 37:111-115 [Vol. III, No. 1]
60. With F. J. MURRAY. On Rings of Operators. *Ann. Math.* 37:116-229 [Vol. III, No. 2]
61. On an Algebraic Generalization of the Quantum Mechanical Formalism (Part I). *Mat. Sborn.* 1:415-484 [Vol. III, No. 9]
62. The Uniqueness of Haar's Measure. *Mat. Sborn.* 1:721-734 [Vol. IV, No. 6]
63. With G. BIRKHOFF. The Logic of Quantum Mechanics. *Ann. Math.* 37:823-843 [Vol. IV, No. 7]
64. Continuous Geometry. *N.A.S. Proc.* 22:92-100 [Vol. IV, No. 8]
65. Examples of Continuous Geometries. *N.A.S. Proc.* 22:101-108 [Vol. IV, No. 9]
66. On Regular Rings. *N.A.S. Proc.* 22:707-713 [Vol. IV, No. 10]

1937

67. With K. KURATOWSKI. On Some Analytic Sets Defined by Transfinite Induction. *Ann. Math.* 38:521-525 [Vol. IV, No. 19]
68. With F. J. MURRAY. On Rings of Operators, II. *Amer. Math. Soc. Trans.* 41:208-248 [Vol. III, No. 3]
69. Some Matrix-Inequalities and Metrization of Matrix-Space. *Tomsk. Univ. Rev.* 1:286-300 [Vol. IV, No. 20]
70. Algebraic Theory of Continuous Geometries. *N.A.S. Proc.* 23:16-22 [Vol. IV, No. 11]
71. Continuous Rings and Their Arithmetics. *N.A.S. Proc.* 23:341-349 [Vol. IV, No. 12]

72. ‡ Über ein ökonomisches Gleichungssystem und eine Verallgemeinerung des Brouwerschen Fixpunktsatzes. *Erg. eines Math. Coll.* Vienna, ed. by K. MENGER, 8:73-83.

1938

73. On Infinite Direct Products. *Compos. Math.* 6:1-77 [Vol. III, No. 6]

1940

74. With I. HALPERIN. On the Transitivity of Perspective Mappings. *Ann. Math.* 41:87-93 [Vol. IV, No. 13]
 75. On Rings of Operators, III. *Ann. Math.* 41:94-161 [Vol. III, No. 4]
 76. With E. WIGNER. Minimally Almost Periodic Groups. *Ann. Math.* 41:746-750 [Vol. IV, No. 21]
 77. ‡ With R. H. KENT. The Estimation of the Probable Error from Successive Differences. Ballistic Research Laboratory Report No. 175, February 14. 19 pp.

1941

78. With R. H. KENT, H. R. BELLINSON and B. I. HART. The Mean Square Successive Difference. *Ann. Math. Stat.* 12:153-162 [Vol. IV, No. 33]
 79. With I. J. SCHOENBERG. Fourier Integrals and Metric Geometry. *Amer. Math. Soc. Trans.* 50:226-251 [Vol. IV, No. 22]
 80. Distribution of the Ratio of the Mean Square Successive Difference to the Variance. *Ann. Math. Stat.* 12:367-395 [Vol. IV, No. 34]
 81. ‡ Shock Waves Started by an Infinitesimally Short Detonation of Given (Positive and Finite) Energy. National Defense Research Council, Div. 8, June 30. AM-9.
 82. Optimum Aiming at an Imperfectly Located Target. Appendix to : *Optimum Spacing of Bombs of Shots in the Presence of Systematic Errors*, by L. S. DEDERICK and R. H. KENT. Ballistic Research Laboratory Report 241, July 3. 26 pp. [Vol. IV, No. 37]

1942

83. A Further Remark Concerning the Distribution of the Ratio of the Mean Square Successive Difference to the Variance. *Ann. Math. Stat.* 13:86-88 [Vol. IV, No. 35]
 84. With P. R. HALMOS. Operator Methods in Classical Mechanics, II. *Ann. Math.* 43:332-350 [Vol. IV, No. 23]
 85. With S. CHANDRASEKHAR. The Statistics of the Gravitational Field Arising from a Random Distribution of Stars, I. *Astrophys. J.* 95:489-531 [Vol. VI, No. 12]
 86. Note to : Tabulation of the Probabilities for the Ratio of the Mean Square Successive Difference to the Variance, by B. I. HART. *Ann. Math. Stat.* 13:207-214 [Vol. IV, No. 36]
 87. Approximative Properties of Matrices of High Finite Order. *Portugaliae Math.* 3:1-62 [Vol. IV, No. 24]
 88. Theory of Detonation Waves. A progress report to April 1, 1942. PB 31090, May 4. 34 pp. [Vol. VI, No. 20]

1943

89. With F. J. MURRAY. On Rings of Operators, IV. *Ann. Math.* 44:716-808 [Vol. III, No. 5]
 90. With S. CHANDRASEKHAR. The Statistics of the Gravitational Field Arising from a Random Distribution of Stars. II. The Speed of Fluctuations ; Dynamical Friction ; Spatial Correlations. *Astrophys. J.* 97:1-27 [Vol. VI, No. 13]
 91. On Some Algebraical Properties of Operator Rings. *Ann. Math.* 44:709-715 [Vol. III, No. 8]
 92. ‡ With R. J. SEEGER. On Oblique Reflection and Collision of Shock Waves. PB 31918, September 20. 3 pp.
 93. Theory of Shock Waves. Progress report to Aug. 31, 1942. PB 32719, January 29. 37 pp. [Vol. VI, No. 19]
 94. *Oblique Reflection of Shocks*. PB 37079, October 12. 75 pp. [Vol. VI, No. 22]

1944

95. Proposal and Analysis of a Numerical Method for the Treatment of Hydrodynamical Shock Problems. OSRD-3617, March 20. 31 pp. [Vol. VI, No. 26]

‡ These items will not be found in the collected works. The Appendix to the Bibliography states the nature of the material omitted.

96. ‡ Introductory Remarks (Sec. I), Theory of the Spinning Detonation (Sec. XII), Theory of the Intermediate Product (Sec. XIII). In : *Report of Informal Technical Conference on the Mechanism of Detonation*. AM-570, April 10. 19 pp.
97. ‡ Riemann Method ; Shock Waves and Discontinuities (one-dimensional), Two-Dimensional Hydrodynamics. Lectures in : *Shock Hydrodynamics and Blast Waves*, by H. A. BETHE, K. FUCHS, J. VON NEUMANN, R. PEIERLS and W. G. PENNEY. Notes by J. O. HIRSCHFELDER. AECD-2860, October 28. 106 pp.

1945

98. A Model of General Economic Equilibrium. *Rev. Econ. Studies* 13:1-9 [Vol. VI, No. 3]
99. Refraction, Intersection and Reflection of Shock Waves. In : *Conference on Supersonic Flow and Shock Waves*, pp. 4-12. AM-1663, July 16 [Vol. VI, No. 23]
100. ‡ With A. H. TAUB. Flying Wind Tunnel Experiments. PB 33263, November 5. 16 pp.

1946

101. With V. BARGMANN and D. MONTGOMERY. Solution of Linear Systems of High Order. Report prepared for Navy BuOrd. Under Contract Nord-9596. 85 pp. [Vol. V, No. 13]
102. With A. W. BURKS and H. H. GOLDSTINE. Preliminary Discussion of the Logical Design of an Electronic Computing Instrument. Part I, Vol. I. Report prepared for U.S. Army Ord. Dept. under Contract W-36-034-ORD-7481. 42 pp. [Vol. V, No. 2]
103. With R. SCHATTEN. The Cross-Space of Linear Transformations. II. *Ann. Math.* 47: 608-630 [Vol. IV, No. 29]
104. With H. H. GOLDSTINE. On the Principles of Large Scale Computing Machines. Unpublished [Vol. V, No. 1]

1947

105. The Mathematician. In : *The Works of the Mind*, ed. by R. B. HEYWOOD (University of Chicago Press), pp. 180-196 [Vol. I, No. 1]
106. With H. H. GOLDSTINE. Numerical Inverting of Matrices of High Order. *Amer. Math. Soc. Bull.* 53:1021-1099 [Vol. V, No. 14]
107. With H. H. GOLDSTINE. Planning and Coding of Problems for an Electronic Computing Instrument. Part II, Vol. I. Report prepared for U.S. Army Ord. Dept. under Contract W-36-034-ORD-7481. 69 pp. [Vol. V, No. 3]
108. With R. D. RICHTMYER. Statistical Methods in Neutron Diffusion. LAMS-551, April 9. 22 pp. [Vol. V, No. 21]
109. The Point Source Solution. Chapt. 2 of *Blast Wave*. LA-2000, August 13, pp. 27-55 [Vol. VI, No. 21]
110. With F. REINES. The Mach Effect and the Height of Burst. Chapt. 10 of *Blast Wave*. LA-2000, August 13, pp. X-11—X-84 [Vol. VI, No. 24]
111. With R. D. RICHTMYER. On the Numerical Solution of Partial Differential Equations of Parabolic Type. LA-657, December 25. 17 pp. [Vol. V, No. 18]

1948

112. With H. H. GOLDSTINE. Planning and Coding of Problems for an Electronic Computing Instrument. Part II, Vol. II. Report prepared for U.S. Army Ord. Dept. under Contract W-36-034-ORD-7481. 68 pp. [Vol. V, No. 4]
113. With H. H. GOLDSTINE. Planning and Coding of Problems for an Electronic Computing Instrument. Part II, Vol. III. Report prepared for U.S. Army Ord. Dept. under Contract W-36-034-ORD-7481. 23 pp. [Vol. V, No. 5]
114. On the Theory of Stationary Detonation Waves. File No. X122, B. R. L., Aberdeen Proving Ground, Md., September 20. 26 pp.‡
115. First Report on the Numerical Calculation of Flow Problems. June 22-July 6. 53 pp. [Vol. V, No. 19]
116. Second Report on the Numerical Calculation of Flow Problems. July 25-August 22. 72 pp. [Vol. V, No. 20]
117. With R. SCHATTEN. The Cross-Space of Linear Transformations. III. *Ann. Math.* 49:557-582 [Vol. IV, No. 30]

1949

118. On Rings of Operators. Reduction Theory. *Ann. Math.* 50:401-485 [Vol. III, No. 7]
119. Recent Theories of Turbulence. (Report made to Office of Naval Research.) Unpublished [Vol. VI, No. 32]

1950

120. With R. D. RICHTMYER. A Method for the Numerical Calculation of Hydrodynamic Shocks. *J. Appl. Phys.* **21**:232-237 [Vol. VI, No. 27]
121. With G. W. BROWN. Solutions of Games by Differential Equations. In : *Contributions to the Theory of Games* (*Ann. Math. Studies* No. 24, Princeton Univ. Press), pp. 73-79 [Vol. VI, No. 4]
122. With N. C. METROPOLIS and G. REITWIESNER. Statistical Treatment of Values of First 2000 Decimal Digits of e and of Pi Calculated on the ENIAC. *Math. Tables and Other Aids to Comp.* **4**:109-111 [Vol. V, No. 22]
123. With J. G. CHARNEY and R. FJORTØFT. Numerical Integration of the Barotropic Vorticity Equation. *Tellus* **2**:237-254 [Vol. VI, No. 29]
124. With I. E. SEGAL. A Theorem on Unitary Representations of Semisimple Lie Groups. *Ann. Math.* **52**:509-517 [Vol. IV, No. 25]

1951

125. Eine Spektraltheorie für allgemeine Operatoren eines unitären Raumes. *Math. Nach.* **4**:258-281 [Vol. IV, No. 26]
126. The Future of High-Speed Computing. *Proc., Comp. Sem.* Dec. 1949, published and copyrighted by I.B.M. (1951), p. 13 [Vol. V, No. 6]
127. Discussion on the Existence and Uniqueness or Multiplicity of Solutions of the Aerodynamical Equations. Chapter 10 of *Problems of Cosmical Aerodynamics*, Proceedings of the Symposium on the Motion of Gaseous Masses of Cosmical Dimensions held at Paris, August 16-19, 1949. Central Air Doc. Office, pp. 75-84 [Vol. VI, No. 25]
128. With H. H. GOLDSTINE. Numerical Inverting of Matrices of High Order, II. *Amer. Math. Soc. Proc.* **2**:188-202 [Vol. V, No. 15]
129. Various Techniques Used in Connection with Random Digits. Chapter 13 of *Proceedings of Symposium on "Monte Carlo Method"* held June-July 1949 in Los Angeles. *J. Res. Nat. Bur. Stand. Appl. Math. Series*, **12**:36-38 [Vol. V, No. 23]
130. The General and Logical Theory of Automata. In : *Cerebral Mechanisms in Behavior—The Hixon Symposium* (September 1948, Pasadena), ed. by L. A. JEFFRESS (John Wiley, New York), pp. 1-31 [Vol. V, No. 9]

1952

131. Discussion Remark Concerning Paper of C. S. SMITH, Grain Shapes and Other Metallurgical Applications of Topology. In : *Metal Interfaces*, Amer. Soc. for Metals, Cleveland, Ohio, pp. 108-110 [Vol. VI, No. 39]

1953

132. A Certain Zero-Sum Two-Person Game Equivalent to the Optimal Assignment Problem. In : *Contributions to the Theory of Games*, Vol. II (*Ann. Math. Studies*, No. 28, Princeton University Press), pp. 5-12 [Vol. VI, No. 5]
133. With D. B. GILLIES and J. P. MAYBERRY. Two Variants of Poker. In : *Contributions to the Theory of Games*, Vol. II (*Ann. Math. Studies*, No. 28, Princeton University Press), pp. 13-50 [Vol. VI, No. 6]
134. With H. H. GOLDSTINE. A Numerical Study of a Conjecture of Kummer. *M.T.A.C.* **VII**, No. 42, pp. 133-134 [Vol. V, No. 24]
135. Communication on the Borel Notes. *Econom.* **21**:124-125 [Vol. VI, No. 2]
136. With E. FERMI. *Taylor Instability at the Boundary of Two Incompressible Liquids*. AECU-2979, Part 2, August 19, pp. 7-13 [Vol. VI, No. 30]

1954

137. With E. P. WIGNER. Significance of Loewner's Theorem in the Quantum Theory of Collisions. *Ann. Math.* **59**:418-433 [Vol. IV, No. 27]
138. The Role of Mathematics in the Sciences and in Society. Address at 4th Conference of Association of Princeton. Graduate Alumni, June 1954, pp. 16-29 [Vol. VI, No. 34]
139. A Numerical Method to Determine Optimum Strategy. *Naval Res. Logistics Quart.* **1**:109-115 [Vol. VI, No. 7]
140. The NORC and Problems in High Speed Computing. Speech at first public showing of I.B.M. Naval Ordnance Research Calculator, December 2, 1954 [Vol. V, No. 7]
141. Entwicklung und Ausnutzung neuerer mathematischer Maschinen. *Arbeitsgemeinschaft für Forschung des Landes Nordrhein-Westfalen*, Vol. 45. Düsseldorf [Vol. V, No. 8]
142. Non-Linear Capacitance or Inductance Switching, Amplifying and Memory Devices. Unpublished. (Basic paper for *Patent 2,815,488*, filed April 28, 1954) [Vol. V, No. 11]

1955

143. Can We Survive Technology ? *Fortune*, June [Vol. VI, No. 36]
144. With A. DEVINATZ and A. E. NUSSBAUM. On the Permutability of Self-Adjoint Operators. *Ann. Math.* 62:199-203 [Vol. IV, No. 28]
145. With BRYANT TUCKERMAN. Continued Fraction Expansion of $2^{1/3}$. *Math. Tables and Other Aids to Computation*, IX:23-24 [Vol. V, No. 25]
146. With H. H. GOLDSTINE. Blast Wave Calculation. *Comm. Pure Appl. Math.* 8:327-353 [Vol. VI, No. 28]
147. Method in the Physical Sciences. In : *The Unity of Knowledge*, ed. by L. LEARY (Doubleday, New York), pp. 157-164 [Vol. VI, No. 35]
148. Defense in Atomic War. (Paper delivered at a symposium in honor of Dr. Robert H. Kent, December 7, 1955.) *The Scientific Bases of Weapons*, pp. 21-23 [Vol. VI, No. 38]
149. Impact of Atomic Energy on the Physical and Chemical Sciences. (Speech at M.I.T. Alumni Day Symposium.) *Tech. Rev.* November, pp. 15-17 [Vol. VI, No. 37]

1956

150. Probabilistic Logics and the Synthesis of Reliable Organisms from Unreliable Components. In : *Automata Studies*, ed. by C. E. SHANNON and J. MCCARTHY (Princeton University Press), pp. 43-98 [Vol. V, No. 10]
151. The Impact of Recent Developments in Science on the Economy and on Economics. (Speech at National Planning Assoc., Washington, D.C. Dec. 12, 1955.) *Looking Ahead* 4:11 [Vol. VI, No. 11]

1958

152. Non-Isomorphism of Certain Continuous Rings (with introduction by I. KAPLANSKY). *Ann. Math.* 67:485-496 [Vol. IV, No. 14]

1959

153. With H. H. GOLDSTINE and F. J. MURRAY. The Jacobi Method for Real Symmetric Matrices. *J. Assoc. Computing Machinery* 6:59-96 [Vol. V, No. 16]
154. With A. BLAIR, N. METROPOLIS, A. H. TAUB and M. TSINGOU. A Study of a Numerical Solution to a Two-Dimensional Hydrodynamical Problem. *Math. Tables and Other Aids to Computation*. XIII:145-184 [Vol. V, No. 17]

Appendix to Bibliography of John von Neumann

The following paragraphs describe the nature of the articles in the bibliography which are omitted from the collected works. The articles are identified by the number appearing in the bibliography.

72. *Über ein ökonomisches Gleichungssystem und eine Verallgemeinerung des Brouwerschen Fixpunktsatzes.*
This paper has been translated and published as item number 98 of the bibliography, which is included in Vol. VI, No. 3.
77. With R. H. KENT. *The Estimation of the Probable Error from Successive Differences.*
The results contained in this report are given in item 78, which is included in Vol. IV, No. 33.
81. *Shock Waves Started by an Infinitesimally Short Detonation of Given (Positive and Finite) Energy.*
Item 109 of the bibliography, which is in Vol. VI, No. 21, is another, more polished version of this report.
92. With R. J. SEEGER. *On Oblique Reflection and Collision of Shock Waves.*
The contents of this short memorandum are contained in item 94 of the bibliography. The latter report is in Vol. VI, No. 22.
96. Introductory Remarks (Sec. I), Theory of the Spinning Detonation (Sec. XII), Theory of the Intermediate Product (Sec. XIII). In : *Report of Informal Technical Conference on the Mechanism of Detonation.*
The remarks made by von Neumann at the conference summarized in this document are in the main contained in item 88 of the bibliography. The latter report is in Vol. VI, No. 20.
97. Riemann Method ; Shock Waves and Discontinuities (one-dimensional) ; Two Dimensional Hydrodynamics. Lectures in : *Shock Hydrodynamics and Blast Waves*, by H. A. BETHE, K. FUCHS, J. VON NEUMANN, R. PEIERLS and W. G. PENNEY. Notes by J. O. HIRSCHFELDER.
This report contains notes of lectures given by von Neumann at Los Alamos Scientific Laboratory. The material in these notes is described in a number of treatises and text books.
100. With A. H. TAUB. *Flying Wind Tunnel Experiments.*
This report describes the first steps of an experimental program which was not carried to completion.
114. *On the Theory of Stationary Detonation Waves.*
Item 88 of the bibliography, which is in Vol. VI, No. 20, contains all the results described in this report and the mathematical derivation of the results.

COMPLETE CONTENTS OF VOLUMES I TO VI

VOLUME I

Logic, Theory of Sets and Quantum Mechanics—The Mathematician. Über die Lage der Nullstellen gewisser Minimumpolynome. Zur Einführung der transfiniten Zahlen. Eine Axiomatisierung der Mengenlehre. Egyenletesen sűrű szamoszorotatok. Zur Prüferschen Theorie der idealen Zahlen. Über die Grundlagen der Quantenmechanik. Zur Theorie der Darstellungen kontinuierlicher Gruppen. Mathematische Begründung der Quantenmechanik. Wahrscheinlichkeitstheoretischer Aufbau der Quantenmechanik. Thermodynamik quantenmechanischer Gesamtheiten. Zur Hilbertschen Beweistheorie. Die Zerlegung eines Intervalles in abzählbar viele kongruente Teilmengen. Ein System algebraisch unabhängiger Zahlen. Über die Definition durch transfinite Induktion und verwandte Fragen der allgemeinen Mengenlehre. Die Axiomatisierung der Mengenlehre. Einige Bemerkungen zur Diracschen Theorie des Drehelektrons. Zur Erklärung einiger Eigenschaften der Spektren aus der Quantenmechanik des Drehelektrons, I. Zur Erklärung einiger Eigenschaften der Spektren aus der Quantenmechanik des Drehelektrons, II. Zur Erklärung einiger Eigenschaften der Spektren aus der Quantenmechanik des Drehelektrons, III. Über eine Widerspruchsfreiheitsfrage der axiomatischen Mengenlehre. Über die analytischen Eigenschaften von Gruppen linearer Transformationen und ihrer Darstellungen. Über merkwürdige diskrete Eigenwerte. Über das Verhalten von Eigenwerten bei adiabatischen Prozessen. Beweis des Ergodensatzes und des H -Theorems in der neuen Mechanik. Zur allgemeinen Theorie des Masses. Zusatz zur Arbeit "Zur allgemeinen Theorie des Masses."

VOLUME II

Operators, Ergodic Theory and Almost Periodic Functions in a Group—Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren. Zur Algebra der Funktionaloperatoren und Theorie der normalen Operatoren. Zur Theorie der unbeschränkten Matrizen. Über einen Hilfssatz der Variationsrechnung. Über Funktionen von Funktionaloperatoren. Algebraische Repräsentanten der Funktionen "bis auf eine Menge vom Masse Null." Die Eindeutigkeit der Schrödingerschen Operatoren. Bemerkungen zu den Ausführungen von Herrn St. Lesniewski über meine Arbeit "Zur Hilbertschen Beweistheorie." Die formalistische Grundlegung der Mathematik. Zum Beweise des Minkowskischen Satzes über Linearformen. Über adjungierte Funktionaloperatoren. Proof of the quasi-ergodic hypothesis. Physical applications of the ergodic hypothesis. Dynamical systems of continuous spectra. Über einen Satz von Herrn M. H. Stone. Einige Sätze Über messbare Abbildungen. Zur Operatorenmethode in der klassischen Mechanik. Zusätze zur Arbeit "Zur Operatorenmethode in der klassischen Mechanik." Die Einführung analytischer Parameter in topologischen Gruppen. A koordináta-mérés pontosságának határai az elektron Dirac-féle elméletében (Über die Grenzen der Koordinatenmessungs-Genauigkeit in der Diracschen Theorie des Elektrons). On an algebraic generalization of the quantum mechanical formalism. Zum Haarschen Mass in topologischen Gruppen. Almost periodic functions in a group. I. The Dirac equation in projective relativity. On complete topological spaces. Almost periodic functions in groups. II. Comparison of Cells (Reviewed by G. Hunt).

VOLUME III

Rings of Operators—On a certain topology for rings of operators. On rings of operators. On rings of operators, II. On rings of operators. III. On rings of operators, IV. On infinite direct products. On rings of operators; Reduction Theory. On some algebraical properties of operator rings. On an algebraic generalization of the quantum mechanical formalism (Part I). Characterization of Factors of Type II_1 (Reviewed by I. Kaplansky).

VOLUME IV

Continuous Geometry and other topics—On compact solutions of operational-differential equations. I. Charakterisierung des Spektrums eines Integraloperators. On normal operators. On inner products in linear, metric spaces. The determination of representative elements in the residual classes of a Boolean algebra. The uniqueness of Haar's measure. The logic of quantum mechanics. Continuous geometry. Examples of continuous geometries. On regular rings. Algebraic theory of continuous geometries. Continuous rings and their arithmetics. On the transitivity of perspective mappings. Non-isomorphism of certain continuous rings (with introduction by I. Kaplansky). Independence of F_∞ of the sequence v (Reviewed by I. Halperin).

Continuous geometries with a transition probability (Reviewed by I. Halperin). Quantum logics. (Strict- and probability-logics) (Reviewed by A. H. Taub). Lattice abelian groups. (Reviewed by G. Birkhoff). On some analytic sets defined by transfinite induction. Some matrix-inequalities and metrization of matrix-space. Minimally almost periodic groups. Fourier integrals and metric geometry. Operator methods in classical mechanics, II. Approximative properties of matrices of high finite order. A theorem on unitary representations of semisimple lie groups. Eine Spektraltheorie für allgemeine Operatoren eines unitären Raumes. Significance of Loewner's theorem in the quantum theory of collisions. On the permutability of self-adjoint operators. The cross-space of linear transformations. II. The cross-space of linear transformations. III. Measure in Functional Spaces (Reviewed by I. Halperin). Representation of certain linear groups by unitary operators in Hilbert space (Reviewed by G. W. Mackey). The mean square successive difference. Distribution of the ratio of the mean square successive difference to the variance. A further remark concerning the distribution of the ratio of the mean square successive difference to the variance. Tabulation of the probabilities for the ratio of the mean square successive difference to the variance. Optimum Aiming at an Imperfectly Located Target.

VOLUME V

Design of Computers, Theory of Automata and Numerical Analysis—On the Principles of Large Scale Computing Machines. Preliminary discussion of the logical design of an electronic computing instrument. Part I, Vol. I. Planning and coding of problems for an electronic computing instrument. Part II, Vol. I. Planning and coding of problems for an electronic computing instrument. Part II, Vol. II. Planning and coding of problems for an electronic computing instrument. Part II, Vol. III. The future of high-speed computing. The NORC and problems in high speed computing. Entwicklung und Ausnutzung neuerer mathematischer Maschinen. The General and Logical Theory of Automata. Probabilistic Logics and the Synthesis of Reliable Organisms from Unreliable Components. Non-linear capacitance or inductance switching, amplifying and memory devices. Notes on the photon-disequilibrium-amplification scheme (Reviewed by J. Bardeen). Solution of linear systems of high order. Numerical inverting of matrices of high order. Numerical inverting of matrices of high order, II. The Jacobi Method for real symmetric matrices. A study of a numerical solution to a two-dimensional hydrodynamical problem. On the numerical solution of partial differential equations of parabolic type. First report on the numerical calculation of flow problems. Second report on the numerical calculation of flow problems. Statistical methods in neutron diffusion. Statistical treatment of values of first 2000 decimal digits of e and of π calculated on the ENIAC. Various techniques used in connection with random digits. A numerical study of a conjecture of Kummer. Continued Fraction Expansion of $2^{1/3}$.

VOLUME VI

Theory of Games, Astrophysics, Hydrodynamics and Meteorology—Zur theorie der Gesellschaftsspiele. Communication on the Borel Notes. A model of general economic equilibrium. Solutions of games by differential equations. A certain zero-sum two-person game equivalent to the optimal assignment problem. Two variants of poker. A numerical method to determine optimum strategy. Discussion of a maximum problem (Reviewed by A. W. Tucker and H. W. Kuhn). Numerical method for determination of value and best strategies of 0-sum, 2-person game with large number of strategies (Reviewed by A. W. Tucker and H. W. Kuhn). Symmetric solution of some general n person games (Reviewed by D. B. Gillies). The impact of recent developments in science on the economy and on economics. The statistics of the gravitational field arising from a random distribution of stars, I. The statistics of the gravitational field arising from a random distribution of stars, II. Static solution of Einstein field equation for perfect fluid with $T_0\rho=0$ (Reviewed by A. H. Taub). On the relativistic gas-degeneracy and the collapsed configuration of stars (Reviewed by A. H. Taub). The point source model (Reviewed by A. H. Taub). The point source solution, assuming a degeneracy of the semi-relativistic type $p=K\rho^{4/3}$ over the entire star (Reviewed by A. H. Taub). Discussion of DeSitter's space and of Dirac's equation in it (Reviewed by A. H. Taub). Theory of shock waves. Theory of detonation waves. The point source solution. Oblique reflection of shocks. Refraction, intersection and reflection of shock waves. The Mach Effect and the Height of Burst. Discussion on the existence and uniqueness or multiplicity of solutions of the aerodynamical equations. Proposal and analysis of a numerical method for the treatment of hydro-dynamical shock problems. A method for the numerical calculation of hydrodynamic shocks. Blast wave calculation. Numerical integration of the barotropic vorticity equation. Taylor instability at the boundary of two incompressible liquids. Taylor instability problem (Reviewed by H. H. Goldstine). Recent theories of turbulence. Description of the conformal mapping method for the integration of partial differential equation systems with $1+2$ independent variables (Reviewed by A. H. Taub). The role of mathematics in the sciences and in society. Method in the physical sciences. Can we survive technology?. Impact of atomic energy on the physical and Chemical Sciences. Defense in Atomic War. Discussion remark concerning paper of C. S. Smith entitled, "Grain shapes and other metallurgical applications of topology."

Volume II

Operators, Ergodic Theory and Almost Periodic Functions in a Group -

Allgemeine Eigenwert theorie

Hermiteischer Funktionaloperatoren. Zur Algebra der Funktionaloperatoren und Theorie der normalen Operatoren. Zur Theorie der unbeschränkten Matrizen. Über einen Hilfssatz der Variationsrechnung. Über Funktionen von Funktionaloperatoren. Algebraische Repräsentanten der Funktionen "bis auf eine Menge vom Masse Null." Die Eindeutigkeit der Schrödingerschen Operatoren. Bemerkungen zu den Ausführungen von Herrn St. Lesniewski über meine Arbeit "Zur Hilbertschen Beweistheorie." Die formalistische Grundlegung der Mathematik. Zum Beweise des Minkowskischen Satzes über Linearformen. Über adjungierte Funktionaloperatoren. Proof of the quasi-ergodic hypothesis. Physical applications of the ergodic hypothesis. Dynamical systems of continuous spectra. Über einen Satz von Herrn M. H. Stone. Einige Sätze Über messbare Abbildungen. Zur Operatorenmethode in der klassischen Mechanik. Zusätze zur Arbeit "Zur Operatorenmethode in der klassischen Mechanik." Die Einführung analytischer Parameter in topologischen Gruppen. A koordináta-mérés pontosságának határai az elektron Dirac-féle elméletében (Über die Grenzen der Koordinatenmessungs-Genauigkeit in der Diracschen Theorie des Elektrons). On an algebraic generalization of the quantum mechanical formalism. Zum Haarschen Mass in topologischen Gruppen. Almost periodic functions in a group. I. The Dirac equation in projective relativity. On complete topological spaces. Almost periodic functions in groups. II. Comparison of Cells (Reviewed by G. Hunt).